

Foundations of Data Science¹

John Hopcroft

Ravindran Kannan

Version 20/11/2013

These notes are a first draft of a book being written by Hopcroft and Kannan and in many places are incomplete. However, the notes are in good enough shape to prepare lectures for a modern theoretical course in computer science. Please do not put solutions to exercises online as it is important for students to work out solutions for themselves rather than copy them from the internet.

Thanks

JEH

¹Copyright 2011. All rights reserved

Contents

1	Introduction	7
2	High-Dimensional Space	10
2.1	Properties of High-Dimensional Space	12
2.2	The Law of Large Numbers	13
2.3	The High-Dimensional Sphere	15
2.3.1	The Sphere and the Cube in High Dimensions	16
2.3.2	Volume and Surface Area of the Unit Sphere	17
2.3.3	The Volume is Near the Equator	20
2.3.4	The Volume is in a Narrow Annulus	23
2.3.5	The Surface Area is Near the Equator	24
2.4	Volumes of Other Solids	25
2.5	Generating Points Uniformly at Random on the Surface of a Sphere	26
2.6	Gaussians in High Dimension	27
2.7	Bounds on Tail Probability	32
2.8	Applications of the tail bound	35
2.9	Random Projection and Johnson-Lindenstrauss Theorem	37
2.10	Bibliographic Notes	40
2.11	Exercises	41
3	Random Graphs	51
3.1	The $G(n, p)$ Model	51
3.1.1	Degree Distribution	52
3.1.2	Existence of Triangles in $G(n, d/n)$	57
3.2	Phase Transitions	59
3.3	The Giant Component	66
3.4	Branching Processes	74
3.5	Cycles and Full Connectivity	80
3.5.1	Emergence of Cycles	80
3.5.2	Full Connectivity	81
3.5.3	Threshold for $\mathbf{O}(\ln n)$ Diameter	83
3.6	Phase Transitions for Increasing Properties	85
3.7	Phase Transitions for CNF-sat	87
3.8	Nonuniform and Growth Models of Random Graphs	91
3.8.1	Nonuniform Models	91
3.8.2	Giant Component in Random Graphs with Given Degree Distribution	92
3.9	Growth Models	93
3.9.1	Growth Model Without Preferential Attachment	93
3.9.2	Growth Model With Preferential Attachment	100
3.10	Small World Graphs	101
3.11	Bibliographic Notes	106
3.12	Exercises	107

4	Singular Value Decomposition (SVD)	117
4.1	Singular Vectors	118
4.2	Singular Value Decomposition (SVD)	122
4.3	Best Rank k Approximations	123
4.4	Power Method for Computing the Singular Value Decomposition	126
4.5	Applications of Singular Value Decomposition	130
4.5.1	Principal Component Analysis	130
4.5.2	Clustering a Mixture of Spherical Gaussians	130
4.5.3	Spectral Decomposition	135
4.5.4	Singular Vectors and Ranking Documents	136
4.5.5	An Application of SVD to a Discrete Optimization Problem	137
4.6	Singular Vectors and Eigenvectors	140
4.7	Bibliographic Notes	141
4.8	Exercises	142
5	Random Walks and Markov Chains	150
5.1	Stationary Distribution	153
5.2	Electrical Networks and Random Walks	155
5.3	Random Walks on Undirected Graphs with Unit Edge Weights	159
5.4	Random Walks in Euclidean Space	166
5.5	The Web as a Markov Chain	170
5.6	Markov Chain Monte Carlo	174
5.6.1	Metropolis-Hasting Algorithm	175
5.6.2	Gibbs Sampling	177
5.7	Areas and Volumes	179
5.8	Convergence of Random Walks on Undirected Graphs	180
5.8.1	Using Normalized Conductance to Prove Convergence	185
5.9	Bibliographic Notes	188
5.10	Exercises	189
6	Learning and VC-dimension	200
6.1	Learning	200
6.2	Linear Separators, the Perceptron Algorithm, and Margins	202
6.3	Nonlinear Separators, Support Vector Machines, and Kernels	207
6.4	Strong and Weak Learning - Boosting	212
6.5	Number of Examples Needed for Prediction: VC-Dimension	214
6.6	Vapnik-Chervonenkis or VC-Dimension	217
6.6.1	Examples of Set Systems and Their VC-Dimension	218
6.6.2	The Shatter Function	221
6.6.3	Shatter Function for Set Systems of Bounded VC-Dimension	222
6.6.4	Intersection Systems	224
6.7	The VC Theorem	224
6.8	Bibliographic Notes	227

6.9	Exercises	229
7	Algorithms for Massive Data Problems	238
7.1	Frequency Moments of Data Streams	238
7.1.1	Number of Distinct Elements in a Data Stream	239
7.1.2	Counting the Number of Occurrences of a Given Element.	243
7.1.3	Counting Frequent Elements	243
7.1.4	The Second Moment	245
7.2	Matrix Algorithms Using Sampling	249
7.2.1	Matrix Multiplication Using Sampling	249
7.2.2	Sketch of a Large Matrix	251
7.3	Sketches of Documents	254
7.4	Exercises	257
8	Clustering	261
8.1	Some Clustering Examples	261
8.2	A k -means Clustering Algorithm	264
8.3	A Greedy Algorithm for k -Center Criterion Clustering	266
8.4	Spectral Clustering	267
8.5	Recursive Clustering Based on Sparse Cuts	273
8.6	Kernel Methods	274
8.7	Agglomerative Clustering	276
8.8	Dense Submatrices and Communities	279
8.9	Flow Methods	282
8.10	Finding a Local Cluster Without Examining the Whole Graph	284
8.11	Axioms for Clustering	290
8.11.1	An Impossibility Result	290
8.11.2	A Satisfiable Set of Axioms	296
8.12	Exercises	298
9	Topic Models, Hidden Markov Process, Graphical Models, and Belief Propagation	302
9.1	Topic Models	302
9.2	Hidden Markov Model	306
9.3	Graphical Models, and Belief Propagation	311
9.4	Bayesian or Belief Networks	311
9.5	Markov Random Fields	312
9.6	Factor Graphs	314
9.7	Tree Algorithms	314
9.8	Message Passing in general Graphs	316
9.9	Graphs with a Single Cycle	318
9.10	Belief Update in Networks with a Single Loop	319
9.11	Maximum Weight Matching	320
9.12	Warning Propagation	324

9.13	Correlation Between Variables	325
9.14	Exercises	329
10	Other Topics	332
10.1	Rankings	332
10.2	Hare System for Voting	334
10.3	Compressed Sensing and Sparse Vectors	335
10.3.1	Unique Reconstruction of a Sparse Vector	336
10.3.2	The Exact Reconstruction Property	338
10.3.3	Restricted Isometry Property	339
10.4	Applications	341
10.4.1	Sparse Vector in Some Coordinate Basis	341
10.4.2	A Representation Cannot be Sparse in Both Time and Frequency Domains	342
10.4.3	Biological	344
10.4.4	Finding Overlapping Cliques or Communities	345
10.4.5	Low Rank Matrices	345
10.5	Exercises	347
11	Appendix	350
11.1	Asymptotic Notation	350
11.2	Sums of Series	351
11.3	Useful Inequalities	355
11.4	Probability	362
11.4.1	Sample Space, Events, Independence	363
11.4.2	Linearity of Expectation	364
11.4.3	Indicator Variables	364
11.4.4	Variance	364
11.4.5	Variance of the Sum of Independent Random Variables	365
11.4.6	The Central Limit Theorem	365
11.4.7	Median	366
11.4.8	Unbiased Estimators	366
11.4.9	Probability Distributions	366
11.4.10	Maximum Likelihood Estimation MLE	369
11.4.11	Tail Bounds	370
11.5	Chernoff Bounds	371
11.6	Eigenvalues and Eigenvectors	375
11.6.1	Eigenvalues and Eigenvectors	375
11.6.2	Symmetric Matrices	377
11.6.3	Extremal Properties of Eigenvalues	378
11.6.4	Eigenvalues of the Sum of Two Symmetric Matrices	380
11.6.5	Norms	382
11.6.6	Important Norms and Their Properties	383

11.6.7	Linear Algebra	385
11.6.8	Distance between subspaces	387
11.7	Generating Functions	388
11.7.1	Generating Functions for Sequences Defined by Recurrence Relationships	389
11.7.2	The Exponential Generating Function and the Moment Generating Function	391
11.8	Miscellaneous	393
11.8.1	Finite Fields	394
11.8.2	Variational Methods	395
11.8.3	Hash Functions	396
11.8.4	Application of Mean Value Theorem	396
11.8.5	Catalan numbers	397
11.8.6	Sperner's Lemma	398
11.8.7	Prüfer	398
11.9	Exercises	399
	Index	407

Foundations of Data Science[†]

John Hopcroft and Ravindran Kannan

20/11/2013

1 Introduction

Computer science as an academic discipline began in the 60's. Emphasis was on programming languages, compilers, operating systems, and the mathematical theory that supported these areas. Courses in theoretical computer science covered finite automata, regular expressions, context free languages, and computability. In the 70's, algorithms was added as an important component of theory. The emphasis was on making computers useful. Today, a fundamental change is taking place and the focus is more on applications. There are many reasons for this change. The merging of computing and communications has played an important role. The enhanced ability to observe, collect and store data in the natural sciences, in commerce, and in other fields calls for a change in our understanding of data and how to handle it in the modern setting. The emergence of the web and social networks, which are by far the largest such structures, presents both opportunities and challenges for theory.

While traditional areas of computer science are still important and highly skilled individuals are needed in these areas, the majority of researchers will be involved with using computers to understand and make usable massive data arising in applications, not just how to make computers useful on specific well-defined problems. With this in mind we have written this book to cover the theory likely to be useful in the next 40 years, just as automata theory, algorithms and related topics gave students an advantage in the last 40 years. One of the major changes is the switch from discrete mathematics to more of an emphasis on probability, statistics, and numerical methods.

Early drafts of the book have been used for both undergraduate and graduate courses. Background material needed for an undergraduate course has been put in the appendix. For this reason, the appendix has homework problems.

This book starts with the treatment of high dimensional geometry. Modern data in diverse fields such as Information Processing, Search, Machine Learning, etc., is often

[†]Copyright 2011. All rights reserved

represented advantageously as vectors with a large number of components. This is so even in cases when the vector representation is not the natural first choice. Our intuition from two or three dimensional space can be surprisingly off the mark when it comes to high dimensional space. Chapter 2 works out the fundamentals needed to understand the differences. The emphasis of the chapter, as well as the book in general, is to get across the mathematical foundations rather than dwell on particular applications that are only briefly described.

The mathematical areas most relevant to dealing with high-dimensional data are matrix algebra and algorithms. We focus on singular value decomposition, a central tool in this area. Chapter 4 gives a from-first-principles description of this. Applications of singular value decomposition include principal component analysis, a widely used technique which we touch upon, as well as modern applications to statistical mixtures of probability densities, discrete optimization, etc., which are described in more detail.

Central to our understanding of large structures, like the web and social networks, is building models to capture essential properties of these structures. The simplest model is that of a random graph formulated by Erdős and Renyi, which we study in detail proving that certain global phenomena, like a giant connected component, arise in such structures with only local choices. We also describe other models of random graphs.

One of the surprises of computer science over the last two decades is that some domain-independent methods have been immensely successful in tackling problems from diverse areas. Machine learning is a striking example. We describe the foundations of machine learning, both learning from given training examples, as well as the theory of Vapnik-Chervonenkis dimension, which tells us how many training examples suffice for learning. Another important domain-independent technique is based on Markov chains. The underlying mathematical theory, as well as the connections to electrical networks, forms the core of our chapter on Markov chains.

The field of algorithms has traditionally assumed that the input data to a problem is presented in random access memory, which the algorithm can repeatedly access. This is not feasible for modern problems. The streaming model and other models have been formulated to better reflect this. In this setting, sampling plays a crucial role and, indeed, we have to sample on the fly. In Chapter 7 we study how to draw good samples efficiently and how to estimate statistical, as well as linear algebra quantities, with such samples.

One of the most important tools in the modern toolkit is clustering, dividing data into groups of similar objects. After describing some of the basic methods for clustering, such as the k-means algorithm, we focus on modern developments in understanding these, as well as newer algorithms. The chapter ends with a study of clustering criteria.

This book also covers graphical models and belief propagation, ranking and voting,

sparse vectors, and compressed sensing. The appendix includes a wealth of background material.

A word about notation in the book. To help the student, we have adopted certain notations, and with a few exceptions, adhered to them. We use lower case letters for scalar variables and functions, bold face lower case for vectors, and upper case letters for matrices. Lower case near the beginning of the alphabet tend to be constants, in the middle of the alphabet, such as i , j , and k , are indices in summations, n and m for integer sizes, and x , y and z for variables. Where the literature traditionally uses a symbol for a quantity, we also used that symbol, even if it meant abandoning our convention. If we have a set of points in some vector space, and work with a subspace, we use n for the number of points, d for the dimension of the space, and k for the dimension of the subspace.

The term "almost surely" means with probability one. We use $\ln n$ for the natural logarithm and $\log n$ for the base two logarithm. If we want base ten, we will use $\log_{10}$. To simplify notation and to make it easier to read we use $E^2(1-x)$ for $(E(1-x))^2$ and $E(1-x)^2$ for $E((1-x)^2)$.

2 High-Dimensional Space

In many applications data is in the form of vectors. In other applications, data is not in the form of vectors, but could be usefully represented by vectors. The *Vector Space Model* [SWY75] is a good example. In the vector space model, a document is represented by a vector, each component of which corresponds to the number of occurrences of a particular term in the document. The English language has on the order of 25,000 words or terms, so each document is represented by a 25,000 dimensional vector. A collection of n documents is represented by a collection of n vectors, one vector per document. The vectors may be arranged as columns of a $25,000 \times n$ matrix. See Figure 2.1. A query is also represented by a vector in the same space. The component of the vector corresponding to a term in the query, specifies the importance of the term to the query. To find documents about cars that are not race cars, a query vector will have a large positive component for the word car and also for the words engine and perhaps door, and a negative component for the words race, betting, etc.

One needs a measure of relevance or similarity of a query to a document. The dot product or cosine of the angle between the two vectors is an often used measure of similarity. To respond to a query, one computes the dot product or the cosine of the angle between the query vector and each document vector and returns the documents with the highest values of these quantities. While it is by no means clear that this approach will do well for the information retrieval problem, many empirical studies have established the effectiveness of this general approach.

The vector space model is useful in ranking or ordering a large collection of documents in decreasing order of importance. For large collections, an approach based on human understanding of each document is not feasible. Instead, an automated procedure is needed that is able to rank documents with those central to the collection ranked highest. Each document is represented as a vector with the vectors forming the columns of a matrix A . The similarity of pairs of documents is defined by the dot product of the vectors. All pairwise similarities are contained in the matrix product $A^T A$. If one assumes that the documents central to the collection are those with high similarity to other documents, then computing $A^T A$ enables one to create a ranking. Define the total similarity of document i to be the sum of the entries in the i^{th} row of $A^T A$ and rank documents by their total similarity. It turns out that with the vector representation on hand, a better way of ranking is to first find the best fit direction. That is, the unit vector $\mathbf{u}$, for which the sum of squared perpendicular distances of all the vectors to $\mathbf{u}$ is minimized. See Figure 2.2. Then, one ranks the vectors according to their dot product with $\mathbf{u}$. The best-fit direction is a well-studied notion in linear algebra. There is elegant theory and efficient algorithms presented in Chapter 4 that facilitate the ranking as well as applications in many other domains.

In the vector space representation of data, properties of vectors such as dot products,

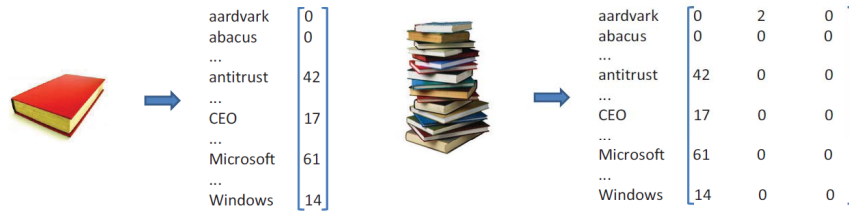


Figure 2.1: A document and its term-document vector along with a collection of documents represented by their term-document vectors.

distance between vectors, and orthogonality, often have natural interpretations and this is what makes the vector representation more important than just a book keeping device. For example, the squared distance between two 0-1 vectors representing links on web pages is the number of web pages linked to by only one of the pages. In Figure 2.3, pages 4 and 5 both have links to pages 1, 3, and 6, but only page 5 has a link to page 2. Thus, the squared distance between the two vectors is one. We have seen that dot products measure similarity. Orthogonality of two nonnegative vectors says that they are disjoint. Thus, if a document collection, e.g., all news articles of a particular year, contained documents on two or more disparate topics, vectors corresponding to documents from different topics would be nearly orthogonal.

The dot product, cosine of the angle, distance, etc., are all measures of similarity or dissimilarity, but there are important mathematical and algorithmic differences between them. The random projection theorem presented in this chapter states that a collection of vectors can be projected to a lower-dimensional space approximately preserving all pairwise distances between vectors. Thus, the nearest neighbors of each vector in the collection can be computed in the projected lower-dimensional space. Such a savings in time is not possible for computing pairwise dot products using a simple projection.

Our aim in this book is to present the reader with the mathematical foundations to deal with high-dimensional data. There are two important parts of this foundation. The first is high-dimensional geometry, along with vectors, matrices, and linear algebra. The second more modern aspect is the combination with probability.

High dimensionality is a common characteristic in many models and for this reason much of this chapter is devoted to the geometry of high-dimensional space, which is quite different from our intuitive understanding of two and three dimensions. We focus first on volumes and surface areas of high-dimensional objects like hyperspheres. We will not present details of any one application, but rather present the fundamental theory useful to many applications.

One reason probability comes in is that many computational problems are hard if our algorithms are required to be efficient on all possible data. In practical situations, domain knowledge often enables the expert to formulate stochastic models of data. In

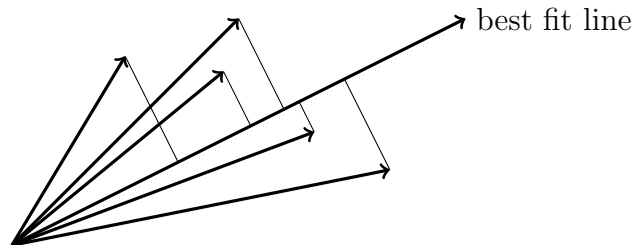


Figure 2.2: The best fit line is the line that minimizes the sum of the squared perpendicular distances.

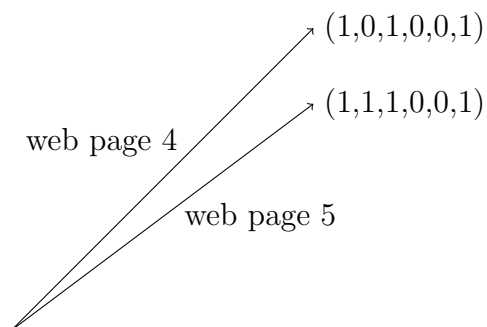


Figure 2.3: Two web pages as vectors. The squared distance between the two vectors is the number of web pages linked to by just one of the two web pages.

customer-product data, a common assumption is that the goods each customer buys are independent of what goods the others buy. One may also assume that the goods a customer buys satisfies a known probability law, like the Gaussian distribution. In keeping with the spirit of the book, we do not discuss specific stochastic models, but present the fundamentals. An important fundamental is the law of large numbers that states that under the assumption of independence of customers, the total consumption of each good is remarkably close to its mean value. The central limit theorem is of a similar flavor. Indeed, it turns out that picking random points from geometric objects like hyperspheres exhibits almost identical properties in high dimensions. One calls this phenomena the “law of large dimensions”. We will establish these geometric properties first before discussing Chernoff bounds and related theorems on aggregates of independent random variables.

2.1 Properties of High-Dimensional Space

Our intuition about space was formed in two and three dimensions and is often misleading in high dimensions. Consider placing 100 points uniformly at random in a unit square. Each coordinate is generated independently and uniformly at random from the interval $[0, 1]$. Select a point and measure the distance to all other points and observe

the distribution of distances. Then increase the dimension and generate the points uniformly at random in a 100-dimensional unit cube. The distribution of distances becomes concentrated about an average distance. The reason is easy to see. Let $\mathbf{x}$ and $\mathbf{y}$ be two such points in d -dimensions. The distance between $\mathbf{x}$ and $\mathbf{y}$ is

$$|\mathbf{x} - \mathbf{y}| = \sqrt{\sum_{i=1}^d (x_i - y_i)^2}.$$

Since $\sum_{i=1}^d (x_i - y_i)^2$ is the summation of a number of independent random variables of bounded variance, by the law of large numbers the distribution of $|\mathbf{x} - \mathbf{y}|^2$ is concentrated about its expected value. Contrast this with the situation where the dimension is two or three and the distribution of distances is spread out.

For another example, consider the difference between picking a point uniformly at random from a unit-radius circle and from a unit-radius sphere in d -dimensions. In d -dimensions the distance from the point to the center of the sphere is very likely to be between $1 - \frac{c}{\sqrt{d}}$ and 1, where c is a constant independent of d . This implies that most of the mass is near the surface of the sphere. Furthermore, the first coordinate, x_1 , of such a point is likely to be between $-\frac{c}{\sqrt{d}}$ and $+\frac{c}{\sqrt{d}}$, which we express by saying that most of the mass is near the equator. The equator perpendicular to the x_1 axis is the set $\{\mathbf{x} | x_1 = 0\}$. We will prove these results in this chapter, but first a review of some probability.

2.2 The Law of Large Numbers

In the previous section, we claimed that points generated at random in high dimensions were all essentially the same distance apart. The reason is that if one averages n independent samples $x_1, x_2, \dots, x_n$ of a random variable x , the result will be close to the expected value of x . Specifically the probability that the average will differ from the expected value by more than ϵ is less than some value $\frac{\sigma^2}{n\epsilon^2}$.

$$\text{Prob} \left(\left| \frac{x_1 + x_2 + \dots + x_n}{n} - E(x) \right| > \epsilon \right) \leq \frac{\sigma^2}{n\epsilon^2}. \quad (2.1)$$

Here the σ^2 in the numerator is the variance of x . The larger the variance of the random variable, the greater the probability that the error will exceed ϵ . The number of points n is in the denominator since the more values that are averaged, the smaller the probability that the difference will exceed ϵ . Similarly the larger ϵ is, the smaller the probability that the difference will exceed ϵ and hence ϵ is in the denominator. Notice that squaring ϵ makes the fraction a dimensionless quantity.

To prove the law of large numbers we use two inequalities. The first is Markov's inequality. One can bound the probability that a nonnegative random variable exceeds a by the expected value of the variable divided by a .

Theorem 2.1 (Markov's inequality) *Let x be a nonnegative random variable. Then for $a > 0$,*

$$\text{Prob}(x \geq a) \leq \frac{E(x)}{a}.$$

Proof: We prove the theorem for continuous random variables. So we use integrals. The same proof works for discrete random variables with sums instead of integrals.

$$\begin{aligned} E(x) &= \int_0^{\infty} xp(x)dx = \int_0^a xp(x)dx + \int_a^{\infty} xp(x)dx \geq \int_a^{\infty} xp(x)dx \\ &\geq \int_a^{\infty} ap(x)dx = a \int_a^{\infty} p(x)dx = a \text{Prob}(x \geq a) \end{aligned}$$

Thus, $\text{Prob}(x \geq a) \leq \frac{E(x)}{a}$. ■

Corollary 2.2 $\text{Prob}(x \geq cE(x)) \leq \frac{1}{c}$

Proof: Substitute $cE(x)$ for a . ■

Markov's inequality bounds the tail of a distribution using only information about the mean. A tighter bound can be obtained by also using the variance.

Theorem 2.3 (Chebyshev's inequality) *Let x be a random variable with mean m and variance σ^2 . Then*

$$\text{Prob}(|x - m| \geq a\sigma) \leq \frac{1}{a^2}.$$

Proof: $\text{Prob}(|x - m| \geq a\sigma) = \text{Prob}((x - m)^2 \geq a^2\sigma^2)$. Note that $(x - m)^2$ is a nonnegative random variable, so Markov's inequality can be applied giving:

$$\text{Prob}((x - m)^2 \geq a^2\sigma^2) \leq \frac{E((x - m)^2)}{a^2\sigma^2} = \frac{\sigma^2}{a^2\sigma^2} = \frac{1}{a^2}.$$

Thus, $\text{Prob}(|x - m| \geq a\sigma) \leq \frac{1}{a^2}$. ■

The law of large numbers follows from Chebyshev's inequality. Recall that $E(x + y) = E(x) + E(y)$, $\sigma^2(cx) = c^2\sigma^2(x)$, $\sigma^2(x - m) = \sigma^2(x)$, and if x and y are independent, then $E(xy) = E(x)E(y)$ and $\sigma^2(x + y) = \sigma^2(x) + \sigma^2(y)$. To prove $\sigma^2(x + y) = \sigma^2(x) + \sigma^2(y)$ when x and y are independent, since $\sigma^2(x - m) = \sigma^2(x)$, one can assume $E(x) = 0$ and $E(y) = 0$. Thus,

$$\begin{aligned} \sigma^2(x + y) &= E((x + y)^2) = E(x^2) + E(y^2) + 2E(xy) \\ &= E(x^2) + E(y^2) + 2E(x)E(y) = \sigma^2(x) + \sigma^2(y). \end{aligned}$$

Replacing $E(xy)$ by $E(x)E(y)$ required independence.

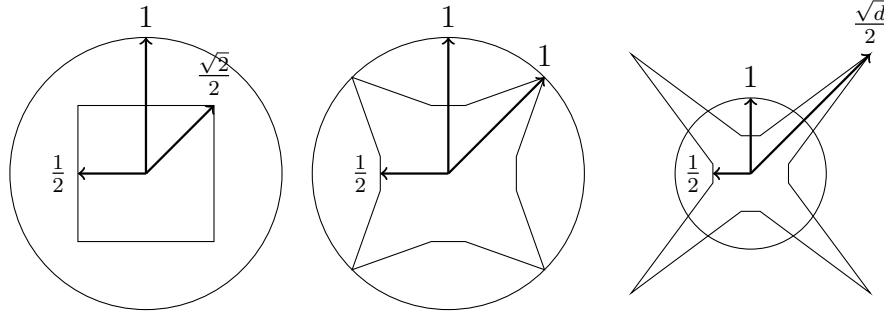


Figure 2.4: Illustration of the relationship between the sphere and the cube in 2, 4, and d -dimensions.

Theorem 2.4 (Law of large numbers) *Let $x_1, x_2, \dots, x_n$ be n samples of a random variable x . Then*

$$\text{Prob} \left(\left| \frac{x_1 + x_2 + \dots + x_n}{n} - E(x) \right| > \epsilon \right) \leq \frac{\sigma^2}{n\epsilon^2}$$

Proof: By Chebychev's inequality

$$\begin{aligned} \text{Prob} \left(\left| \frac{x_1 + x_2 + \dots + x_n}{n} - E(x) \right| > \epsilon \right) &\leq \frac{\sigma^2 \left(\frac{x_1 + x_2 + \dots + x_n}{n} \right)}{\epsilon^2} \\ &\leq \frac{1}{n^2 \epsilon^2} \sigma^2(x_1 + x_2 + \dots + x_n) \\ &\leq \frac{1}{n^2 \epsilon^2} (\sigma^2(x_1) + \sigma^2(x_2) + \dots + \sigma^2(x_n)) \\ &\leq \frac{\sigma^2(x)}{n\epsilon^2}. \end{aligned}$$

■

The law of large numbers bounds the difference of the sample average and the expected value. Note that the size of the sample for a given error bound is independent of the size of the population class. In the limit, when the sample size goes to infinity, the central limit theorem says that the distribution of the sample average is Gaussian provided the random variable has finite variance. Later, we will consider random variables that are the sum of random variables. That is, $x = x_1 + x_2 + \dots + x_n$. Chernoff bounds will tell us about the probability of x differing from its expected value. We will delay this until Section 11.5.

2.3 The High-Dimensional Sphere

One of the interesting facts about a unit-radius sphere in high dimensions is that as the dimension increases, the volume of the sphere goes to zero. This has important

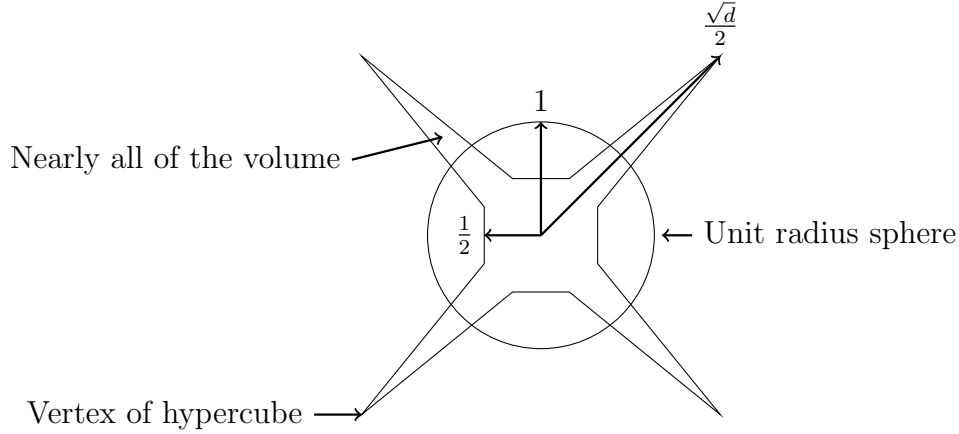


Figure 2.5: Conceptual drawing of a sphere and a cube.

implications. Also, the volume of a high-dimensional sphere is essentially all contained in a thin slice at the equator and simultaneously in a narrow annulus at the surface. There is essentially no interior volume. Similarly, the surface area is essentially all at the equator. These facts, which are contrary to our two or three-dimensional intuition, will be proved by integration.

2.3.1 The Sphere and the Cube in High Dimensions

Consider the difference between the volume of a cube with unit-length sides and the volume of a unit-radius sphere as the dimension d of the space increases. As the dimension of the cube increases, its volume is always one and the maximum possible distance between two points grows as $\sqrt{d}$. In contrast, as the dimension of a unit-radius sphere increases, its volume goes to zero and the maximum possible distance between two points stays at two.

For $d=2$, the unit square centered at the origin lies completely inside the unit-radius circle. The distance from the origin to a vertex of the square is

$$\sqrt{\left(\frac{1}{2}\right)^2 + \left(\frac{1}{2}\right)^2} = \frac{\sqrt{2}}{2} \cong 0.707.$$

Here, the square lies inside the circle. At $d=4$, the distance from the origin to a vertex of a unit cube centered at the origin is

$$\sqrt{\left(\frac{1}{2}\right)^2 + \left(\frac{1}{2}\right)^2 + \left(\frac{1}{2}\right)^2 + \left(\frac{1}{2}\right)^2} = 1.$$

Thus, the vertex lies on the surface of the unit 4-sphere centered at the origin. As the dimension d increases, the distance from the origin to a vertex of the cube increases as $\frac{\sqrt{d}}{2}$, and for large d , the vertices of the cube lie far outside the unit radius sphere. Figure 2.5 illustrates conceptually a cube and a sphere. The vertices of the cube are at distance

$\frac{\sqrt{d}}{2}$ from the origin and for large d lie outside the unit sphere. On the other hand, the mid point of each face of the cube is only distance $1/2$ from the origin and thus is inside the sphere. For large d , almost all the volume of the cube is located outside the sphere.

2.3.2 Volume and Surface Area of the Unit Sphere

For fixed dimension d , the volume of a sphere is a function of its radius and grows as r^d . For fixed radius, the volume of a sphere is a function of the dimension of the space. What is interesting is that the volume of a unit sphere goes to zero as the dimension of the sphere increases.

To calculate the volume of a unit-radius sphere, one can integrate in either Cartesian or polar coordinates. In Cartesian coordinates the volume of a unit sphere is given by

$$V(d) = \int_{x_1=-1}^{x_1=1} \int_{x_2=-\sqrt{1-x_1^2}}^{x_2=\sqrt{1-x_1^2}} \cdots \int_{x_d=-\sqrt{1-x_1^2-\cdots-x_{d-1}^2}}^{x_d=\sqrt{1-x_1^2-\cdots-x_{d-1}^2}} dx_d \cdots dx_2 dx_1.$$

Since the limits of the integrals are complicated, it is easier to integrate using polar coordinates. In polar coordinates, $V(d)$ is given by

$$V(d) = \int_{S^d} \int_{r=0}^1 r^{d-1} dr d\Omega.$$

Here, $d\Omega$ is the surface area of the infinitesimal piece of the solid angle S^d of the unit sphere. See Figure 2.6. The convex hull of the $d\Omega$ piece and the origin form a cone. At radius r , the surface area of the top of the cone is $r^{d-1}d\Omega$ since the surface area is $d-1$ dimensional and each dimension scales by r . The volume of the infinitesimal piece is base times height, and since the surface of the sphere is perpendicular to the radial direction at each point, the height is dr giving the above integral.

Since the variables Ω and r do not interact,

$$V(d) = \int_{S^d} d\Omega \int_{r=0}^1 r^{d-1} dr = \frac{1}{d} \int_{S^d} d\Omega = \frac{A(d)}{d}$$

where $A(d)$ is the surface area of a d -dimensional unit-radius sphere. The question remains, how to determine the surface area $A(d) = \int_{S^d} d\Omega$.

Consider a different integral

$$I(d) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \cdots \int_{-\infty}^{\infty} e^{-(x_1^2+x_2^2+\cdots+x_d^2)} dx_d \cdots dx_2 dx_1.$$

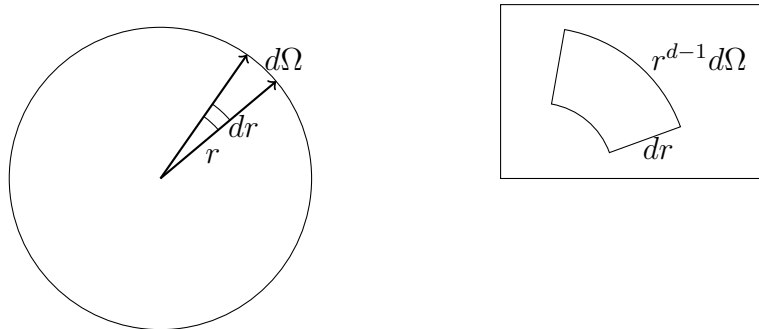


Figure 2.6: Infinitesimal volume in a d -dimensional sphere of unit radius.

Including the exponential allows integration to infinity rather than stopping at the surface of the sphere. Thus, $I(d)$ can be computed by integrating in both Cartesian and polar coordinates. Integrating in polar coordinates will relate $I(d)$ to the surface area $A(d)$. Equating the two results for $I(d)$ allows one to solve for $A(d)$.

First, calculate $I(d)$ by integration in Cartesian coordinates.

$$I(d) = \left[\int_{-\infty}^{\infty} e^{-x^2} dx \right]^d = (\sqrt{\pi})^d = \pi^{\frac{d}{2}}.$$

Here, we have used the fact that $\int_{-\infty}^{\infty} e^{-x^2} dx = \sqrt{\pi}$. For a proof of this, see Section 11.2 of the appendix. Next, calculate $I(d)$ by integrating in polar coordinates. The volume of the differential element is $r^{d-1} d\Omega dr$. Thus,

$$I(d) = \int_{S^d} d\Omega \int_0^{\infty} e^{-r^2} r^{d-1} dr.$$

The integral $\int_{S^d} d\Omega$ is the integral over the entire solid angle and gives the surface area,

$A(d)$, of a unit sphere. Thus, $I(d) = A(d) \int_0^{\infty} e^{-r^2} r^{d-1} dr$. Evaluating the remaining integral gives

$$\int_0^{\infty} e^{-r^2} r^{d-1} dr = \frac{1}{2} \int_0^{\infty} e^{-t} t^{\frac{d}{2}-1} dt = \frac{1}{2} \Gamma\left(\frac{d}{2}\right)$$

and hence, $I(d) = A(d) \frac{1}{2} \Gamma\left(\frac{d}{2}\right)$ where the gamma function $\Gamma(x)$ is a generalization of the factorial function for noninteger values of x . $\Gamma(x) = (x-1)\Gamma(x-1)$, $\Gamma(1) = \Gamma(2) = 1$, and $\Gamma\left(\frac{1}{2}\right) = \sqrt{\pi}$. For integer x , $\Gamma(x) = (x-1)!$.

Cartesian coordinates

$$\underbrace{V(d) = \int \int \cdots \int dx_d \cdots dx_1}_{\text{too hard because of limits}}$$

$$\underbrace{I(d) = \left[\int_{-\infty}^{\infty} e^{-x^2} dx \right]^d}_{\text{evaluate } I(d) \text{ instead}} = \pi^{\frac{d}{2}}$$

Polar coordinates

$\Updownarrow$ equate and solve for $A(d)$

$$V(d) = \int_{S^d} d\Omega \int_{r=0}^1 r^{d-1} dr = \frac{A(d)}{d}$$

$$I(d) = \int_{S^d} d\Omega \int_0^{\infty} e^{-r^2} r^{d-1} dr = \frac{A(d)}{d}$$

$\Longleftarrow$

substitute value of $A(d)$
into formula for $V(d)$

Equate integrals for $I(d)$ in Cartesian and polar coordinates and solve for $A(d)$. Substitute $A(d)$ into the formula for volume of the sphere obtained by integrating in polar coordinates. This gives the result for $V(d)$.

Figure 2.7: Strategy for calculating the volume of a d -dimensional sphere.

Combining $I(d) = \pi^{\frac{d}{2}}$ with $I(d) = A(d) \frac{1}{2} \Gamma\left(\frac{d}{2}\right)$ yields

$$A(d) = \frac{\pi^{\frac{d}{2}}}{\frac{1}{2} \Gamma\left(\frac{d}{2}\right)}$$

establishing the following lemma.

Lemma 2.5 *The surface area $A(d)$ and the volume $V(d)$ of a unit-radius sphere in d dimensions are given by*

$$A(d) = \frac{2\pi^{\frac{d}{2}}}{\Gamma\left(\frac{d}{2}\right)} \quad \text{and} \quad V(d) = \frac{2}{d} \frac{\pi^{\frac{d}{2}}}{\Gamma\left(\frac{d}{2}\right)}.$$

To check the formula for the volume of a unit sphere, note that $V(2) = \pi$ and $V(3) = \frac{2}{3} \frac{\pi^{\frac{3}{2}}}{\Gamma\left(\frac{3}{2}\right)} = \frac{4}{3}\pi$, which are the correct volumes for the unit spheres in two and three dimensions. To check the formula for the surface area of a unit sphere, note that $A(2) = 2\pi$ and $A(3) = \frac{2\pi^{\frac{3}{2}}}{\frac{1}{2}\sqrt{\pi}} = 4\pi$, which are the correct surface areas for the unit sphere in two and three dimensions. Note that $\pi^{\frac{d}{2}}$ is an exponential in $\frac{d}{2}$ and $\Gamma\left(\frac{d}{2}\right)$ grows as the factorial of $\frac{d}{2}$. This implies that $\lim_{d \rightarrow \infty} V(d) = 0$, as claimed.

The volume of a d -dimensional sphere of radius r grows as r^d . This follows since the unit sphere can be mapped to a sphere of radius r by the linear transformation specified by a diagonal matrix with diagonal elements r . The determinant of this matrix is r^d . See Section 2.4. Since the surface area is the derivative of the volume, the surface area grows as r^{d-1} . See last paragraph of Section 2.3.5.

The proof of Lemma 2.5 illustrates the relationship between the surface area of the sphere and the Gaussian probability density

$$\frac{1}{\sqrt{2\pi}} e^{-(x_1+x_2+\dots+x_d)^2/2}.$$

This relationship is an important one and will be used several times in this chapter.

2.3.3 The Volume is Near the Equator

Consider a high-dimensional unit-radius sphere and fix the North Pole on the x_1 axis at $x_1 = 1$. Divide the sphere in half by intersecting it with the plane $x_1 = 0$. The intersection of the plane with the sphere forms a region of one lower dimension, namely $\{\mathbf{x} \mid |\mathbf{x}| \leq 1, x_1 = 0\}$, called the equator. The intersection is a sphere of dimension $d - 1$ and has volume $V(d - 1)$. In three dimensions this region is a circle, in four dimensions the region is a 3-dimensional sphere, etc. In our terminology, a circle is a 2-dimensional sphere and its volume is what one usually refers to as the area of a circle. The surface area of the 2-dimensional sphere is what one usually refers to as the circumference of a circle.

It turns out that essentially all of the volume of the upper hemisphere lies between the plane $x_1 = 0$ and a parallel plane, $x_1 = \varepsilon$, that is slightly higher. For what value of ε does essentially all the volume lie between $x_1 = 0$ and $x_1 = \varepsilon$? The answer depends on the dimension. For dimension d , it is $O(\frac{1}{\sqrt{d-1}})$. To see this, compute the ratio of the volume above the slice lying between $x_1 = 0$ and $x_1 = \varepsilon$ and the volume of the entire upper hemisphere. Actually we compute the ratio of an upper bound on the volume above the slice and a lower bound on the volume of the entire hemisphere and show that this ratio is very small when ε is $O(\frac{1}{\sqrt{d-1}})$.

$$\frac{\text{Volume above slice}}{\text{Volume upper hemisphere}} \leq \frac{\text{Upper bound volume above slice}}{\text{Lower bound volume upper hemisphere}}$$

Let $T = \{\mathbf{x} \mid |\mathbf{x}| \leq 1, x_1 \geq \varepsilon\}$ be the portion of the sphere above the slice. To calculate the volume of T , integrate over x_1 from ε to 1. The incremental volume is a disk of width dx_1 whose face is a $(d-1)$ -dimensional sphere of radius $\sqrt{1 - x_1^2}$. See Figure 2.8. Therefore, the surface area of the disk is

$$(1 - x_1^2)^{\frac{d-1}{2}} V(d - 1)$$

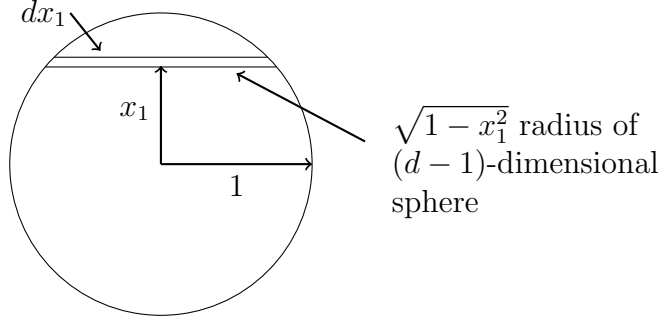


Figure 2.8: The volume of a cross-sectional slab of a d -dimensional sphere.

and

$$\text{Volume}(T) = \int_{\varepsilon}^1 (1 - x_1^2)^{\frac{d-1}{2}} V(d-1) dx_1 = V(d-1) \int_{\varepsilon}^1 (1 - x_1^2)^{\frac{d-1}{2}} dx_1.$$

Note that $V(d)$ denotes the volume of the d -dimensional unit sphere. For the volume of other sets such as the set T , we use the notation $\text{Volume}(T)$ for the volume.

The above integral is difficult to evaluate, so we use some approximations. First, we use the inequality $1 + x \leq e^x$ for all real x and change the upper bound on the integral to infinity. Since x_1 is always greater than ε over the region of integration, we can insert x_1/ε in the integral. This gives

$$\begin{aligned} \text{Volume}(T) &\leq V(d-1) \int_{\varepsilon}^{\infty} e^{-\frac{d-1}{2}x_1^2} dx_1 \\ &\leq V(d-1) \int_{\varepsilon}^{\infty} \frac{x_1}{\varepsilon} e^{-\frac{d-1}{2}x_1^2} dx_1. \end{aligned}$$

Now, $\int x_1 e^{-\frac{d-1}{2}x_1^2} dx_1 = -\frac{1}{d-1} e^{-\frac{d-1}{2}x_1^2}$ and, hence,

$$\text{Volume}(T) \leq \frac{1}{\varepsilon(d-1)} e^{-\frac{d-1}{2}\varepsilon^2} V(d-1). \quad (2.2)$$

The actual volume of the upper hemisphere is exactly $\frac{1}{2}V(d)$. However, we want the volume in terms of $V(d-1)$ instead of $V(d)$ so we can cancel the $V(d-1)$ in the upper bound of the volume above the slice. We do this by calculating a lower bound on the volume of the entire upper hemisphere. Clearly, the volume of the upper hemisphere is at least the volume between the slabs $x_1 = 0$ and $x_1 = \frac{1}{\sqrt{d-1}}$, which is at least the volume of the cylinder of radius $\sqrt{1 - \frac{1}{d-1}}$ and height $\frac{1}{\sqrt{d-1}}$. The volume of the cylinder is $1/\sqrt{d-1}$

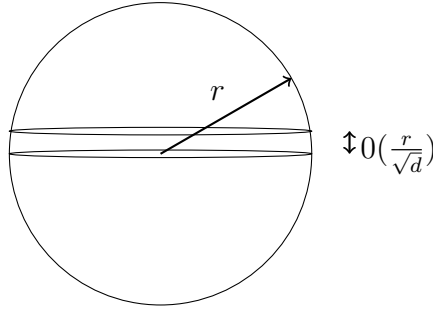


Figure 2.9: Most of the volume of the d -dimensional sphere of radius r is within distance $O(\frac{r}{\sqrt{d}})$ of the equator.

times the $d - 1$ -dimensional volume of the disk $R = \left\{ \mathbf{x} \mid |\mathbf{x}| \leq 1; x_1 = \frac{1}{\sqrt{d-1}} \right\}$. Now R is a $d - 1$ -dimensional sphere of radius $\sqrt{1 - \frac{1}{d-1}}$ and so its volume is

$$\text{Volume}(R) = V(d-1) \left(1 - \frac{1}{d-1} \right)^{(d-1)/2}.$$

Using $(1 - x)^a \geq 1 - ax$

$$\text{Volume}(R) \geq V(d-1) \left(1 - \frac{1}{d-1} \frac{d-1}{2} \right) = \frac{1}{2} V(d-1).$$

Thus, the volume of the upper hemisphere is at least $\frac{1}{2\sqrt{d-1}} V(d-1)$.

The fraction of the volume above the plane $x_1 = \varepsilon$ is upper bounded by the ratio of the upper bound on the volume of the hemisphere above the plane $x_1 = \varepsilon$ to the lower bound on the total volume. This ratio is $\frac{2}{\varepsilon\sqrt{d-1}} e^{-\frac{d-1}{2}\varepsilon^2}$ which leads to the following lemma.

Lemma 2.6 *For any $c > 0$, the fraction of the volume of the unit hemisphere above the plane $x_1 = \frac{c}{\sqrt{d-1}}$ is less than $\frac{2}{c} e^{-c^2/2}$.*

Proof: Substitute $\frac{c}{\sqrt{d-1}}$ for ε in the above. ■

For a large constant c , $\frac{2}{c} e^{-c^2/2}$ is small. The important item to remember is that most of the volume of the d -dimensional unit sphere lies within distance $O(1/\sqrt{d})$ of the equator. If the sphere is of radius r , then the upper bound on the volume above $x_1 = \varepsilon$ increases by r^{d+1} and the lower bound on the volume of the upper hemisphere increases by r^d , which results in an upper bound on the fraction above the plane $x_1 = \varepsilon$ of

$$\frac{2r}{\varepsilon\sqrt{d-1}} e^{-\frac{d-1}{2} \frac{\varepsilon^2}{r^2}}.$$

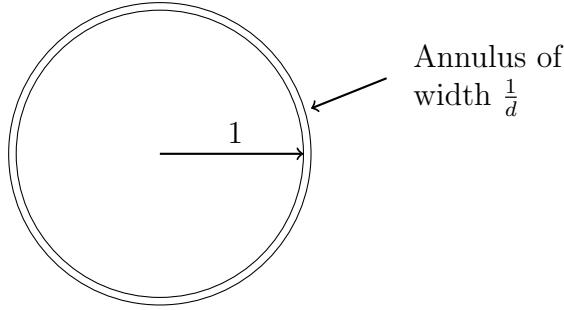


Figure 2.10: Most of the volume of the d -dimensional sphere of radius r is contained in an annulus of width $O(r/d)$ near the boundary.

Substituting $\frac{cr}{\sqrt{d-1}}$ for ϵ , results in a bound of $\frac{2}{c}e^{-\frac{c^2}{2}}$. Thus, most of the volume of a radius r sphere lies within distance $O(\frac{r}{\sqrt{d}})$ of the equator as shown in Figure 2.9.

For $c \geq 2$, the fraction of the volume of the hemisphere above $x_1 = \frac{c}{\sqrt{d-1}}$ is less than $e^{-2} \approx 0.14$ and for $c \geq 4$ the fraction is less than $\frac{1}{2}e^{-8} \approx 3 \times 10^{-4}$. Essentially all the volume of the sphere lies in a narrow band at the equator.

Note that we selected a unit vector in the x_1 direction and defined the equator to be the intersection of the sphere with a $(d-1)$ -dimensional plane perpendicular to the unit vector. However, we could have selected an arbitrary point on the surface of the sphere and considered the vector from the center of the sphere to that point and defined the equator using the plane through the center perpendicular to this arbitrary vector. Essentially all the volume of the sphere lies in a narrow band about this equator also.

2.3.4 The Volume is in a Narrow Annulus

The ratio of the volume of a sphere of radius $1 - \epsilon$ to the volume of a unit sphere in d -dimensions is

$$\frac{(1 - \epsilon)^d V(d)}{V(d)} = (1 - \epsilon)^d,$$

and thus goes to zero as d goes to infinity when ϵ is a fixed constant. In high dimensions, all of the volume of the sphere is concentrated in a narrow annulus at the surface.

Since, $(1 - \epsilon)^d \leq e^{-\epsilon d}$, if $\epsilon = \frac{c}{d}$, for a large constant c , all but e^{-c} of the volume of the sphere is contained in a thin annulus of width c/d . The important item to remember is that most of the volume of the d -dimensional unit sphere is contained in an annulus of width $O(1/d)$ near the boundary. If the sphere is of radius r , then for sufficiently large d , the volume is contained in an annulus of width $O(\frac{r}{d})$.

2.3.5 The Surface Area is Near the Equator

Just as a 2-dimensional circle has an area and a circumference and a 3-dimensional sphere has a volume and a surface area, a d -dimensional sphere has a volume and a surface area. The surface of the sphere is the set $\{\mathbf{x} \mid |\mathbf{x}| = 1\}$. The surface of the equator is the set $\{\mathbf{x} \mid |\mathbf{x}| = 1, x_1 = 0\}$ and it is the surface of a sphere of one lower dimension, i.e., for a 3-dimensional sphere, it is the circumference of a circle. Just as with volume, essentially all the surface area of a high-dimensional sphere is near the equator. To see this, we use an analogous argument to that used for volume.

First, upper bound the surface area of the sphere above $x_1 = \varepsilon$. Let $S = \{\mathbf{x} \mid |\mathbf{x}| = 1, x_1 \geq \varepsilon\}$. To calculate the surface area S of the sphere above $x_1 = \varepsilon$, integrate x_1 from ε to 1. The incremental surface unit will be a band of width dx_1 whose edge is the surface area of a $(d-1)$ -dimensional sphere of radius depending on x_1 . The radius of the band is $\sqrt{1-x_1^2}$ and therefore, the surface area of the $(d-1)$ -dimensional sphere is

$$A(d-1) (1-x_1^2)^{\frac{d-2}{2}}$$

where $A(d-1)$ is the surface area of a unit sphere of dimension $d-1$. The slice is not a cylinder since when x_1 increases by dx_1 , the radius r decreases by dr . Thus,

$$A(S) = A(d-1) \int_{\varepsilon}^1 (1-x_1^2)^{\frac{d-2}{2}} ds$$

where $ds^2 = dr^2 + dx_1^2$. Since $r = \sqrt{1-x_1^2}$, $dr = \frac{-x_1}{\sqrt{1-x_1^2}} dx_1$ and hence

$$ds^2 = \left(\frac{x_1^2}{1-x_1^2} + 1 \right) dx_1^2 = \frac{1}{1-x_1^2} dx_1^2$$

and $ds = \frac{1}{\sqrt{1-x_1^2}} dx_1$. Thus,

$$A(S) = A(d-1) \int_{\varepsilon}^1 (1-x_1^2)^{\frac{d-3}{2}} dx_1.$$

The above integral is difficult to integrate and the same approximations, as in the earlier section on volume, lead to the bound

$$A(S) \leq \frac{1}{\varepsilon^{d-3}} e^{-\frac{d-3}{2}\varepsilon^2} A(d-1). \quad (2.3)$$

Next, lower bound the surface area of the entire upper hemisphere. Clearly, the surface area of the upper hemisphere is greater than the surface area of the side of a d -dimensional cylinder of height $\frac{1}{\sqrt{d-2}}$ and radius $\sqrt{1-\frac{1}{d-2}}$. The surface area of the cylinder is $\frac{1}{\sqrt{d-2}}$ times the circumference area of the d -dimensional cylinder of radius $\sqrt{1-\frac{1}{d-2}}$ which is

$A(d-1)(1 - \frac{1}{d-2})^{\frac{d-2}{2}}$. Using $(1-x)^a \geq 1-ax$, the surface area of the hemisphere is at least

$$\begin{aligned} \frac{1}{\sqrt{d-2}}(1 - \frac{1}{d-2})^{\frac{d-2}{2}} A(d-1) &\geq \frac{1}{\sqrt{d-2}}(1 - \frac{d-2}{2} \frac{1}{d-2}) A(d-1) \\ &\geq \frac{1}{2\sqrt{d-2}} A(d-1). \end{aligned} \quad (2.4)$$

Comparing the upper bound on the surface area of S in (2.3) with the lower bound on the surface area of the hemisphere in (2.4), we see that the surface area above the band $\{\mathbf{x} \mid |\mathbf{x}| = 1, 0 \leq x_1 \leq \varepsilon\}$ is less than $\frac{4}{\varepsilon\sqrt{d-3}} e^{-\frac{d-3}{2}\varepsilon^2}$ of the total surface area.

Lemma 2.7 *For any $c > 0$, the fraction of the surface area above the plane $x_1 = \frac{c}{\sqrt{d-2}}$ is less than or equal to $\frac{4}{c} e^{-\frac{c^2}{2}}$.*

Proof: Substitute $\frac{c}{\sqrt{d-2}}$ for ε in the above. ■

We conclude this section by relating the surface area and volume of a d -dimensional sphere. So far, we have considered unit-radius spheres of dimension d . Now fix the dimension d and vary the radius r . Let $V(d, r)$ denote the volume and let $A(d, r)$ denote the surface area of a d -dimensional sphere of radius r . Then,

$$V(d, r) = \int_{x=0}^r A(d, x) dx.$$

Thus, it follows that the surface area is the derivative of the volume with respect to the radius. In two dimensions, the volume of a circle is πr^2 and the circumference is $2\pi r$. In three dimensions, the volume of a sphere is $\frac{4}{3}\pi r^3$ and the surface area is $4\pi r^2$.

2.4 Volumes of Other Solids

There are very few high-dimensional solids for which there are closed-form formulae for the volume. The volume of the rectangular solid

$$R = \{\mathbf{x} \mid l_1 \leq x_1 \leq u_1, l_2 \leq x_2 \leq u_2, \dots, l_d \leq x_d \leq u_d\}$$

is the product of the lengths of its sides. Namely, it is $\prod_{i=1}^d (u_i - l_i)$.

A parallelepiped is a solid described by

$$P = \{\mathbf{x} \mid \mathbf{l} \leq A\mathbf{x} \leq \mathbf{u}\}$$

where A is an invertible $d \times d$ matrix, and $\mathbf{l}$ and $\mathbf{u}$ are lower and upper bound vectors, respectively. The statements $\mathbf{l} \leq A\mathbf{x}$ and $A\mathbf{x} \leq \mathbf{u}$ are to be interpreted row by row asserting $2d$ inequalities. A parallelepiped is a generalization of a parallelogram. It is

easy to see that P is the image under an invertible linear transformation of a rectangular solid. Let

$$R = \{\mathbf{y} \mid \mathbf{l} \leq \mathbf{y} \leq \mathbf{u}\}.$$

The map $\mathbf{x} = A^{-1}\mathbf{y}$ maps R to P . This implies that

$$\text{Volume}(P) = |\text{Det}(A^{-1})| \text{Volume}(R).$$

Simplices, which are generalizations of triangles, are another class of solids for which volumes can be easily calculated. Consider the triangle in the plane with vertices $\{(0,0), (1,0), (1,1)\}$, which can be described as $\{(x,y) \mid 0 \leq y \leq x \leq 1\}$. Its area is $1/2$ because two such right triangles can be combined to form the unit square. The generalization is the simplex in d -space with $d+1$ vertices,

$$\{(0,0,\dots,0), (1,0,0,\dots,0), (1,1,0,0,\dots,0), \dots, (1,1,\dots,1)\},$$

which is the set

$$S = \{\mathbf{x} \mid 1 \geq x_1 \geq x_2 \geq \dots \geq x_d \geq 0\}.$$

How many copies of this simplex exactly fit into the unit square, $\{\mathbf{x} \mid 0 \leq x_i \leq 1\}$? Every point in the square has some ordering of its coordinates. Since there are $d!$ orderings, exactly $d!$ simplices fit into the unit square. Thus, the volume of each simplex is $1/d!$. Now consider the right angle simplex R whose vertices are the d unit vectors $(1,0,0,\dots,0), (0,1,0,\dots,0), \dots, (0,0,0,\dots,0,1)$ and the origin. A vector $\mathbf{y}$ in R is mapped to an $\mathbf{x}$ in S by the mapping: $x_d = y_d$; $x_{d-1} = y_d + y_{d-1}$; $\dots$; $x_1 = y_1 + y_2 + \dots + y_d$. This is an invertible transformation with determinant one, so the volume of R is also $1/d!$.

A general simplex is obtained by a translation, adding the same vector to every point, followed by an invertible linear transformation on the right simplex. Convince yourself that in the plane every triangle is the image under a translation plus an invertible linear transformation of the right triangle. As in the case of parallelepipeds, applying a linear transformation A multiplies the volume by the determinant of A . Translation does not change the volume. Thus, if the vertices of a simplex T are $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{d+1}$, then translating the simplex by $-\mathbf{v}_{d+1}$ results in vertices $\mathbf{v}_1 - \mathbf{v}_{d+1}, \mathbf{v}_2 - \mathbf{v}_{d+1}, \dots, \mathbf{v}_d - \mathbf{v}_{d+1}, \mathbf{0}$. Let A be the $d \times d$ matrix with columns $\mathbf{v}_1 - \mathbf{v}_{d+1}, \mathbf{v}_2 - \mathbf{v}_{d+1}, \dots, \mathbf{v}_d - \mathbf{v}_{d+1}$. Then, $A^{-1}T = R$ and $AR = T$ where R is the right angle simplex. Thus, the volume of T is $\frac{1}{d!}|\text{Det}(A)|$.

2.5 Generating Points Uniformly at Random on the Surface of a Sphere

Consider generating points uniformly at random on the surface of a unit-radius sphere. First, consider the 2-dimensional version of generating points on the circumference of a unit-radius circle by the following method. Independently generate each coordinate uniformly at random from the interval $[-1, 1]$. This produces points distributed over a square

that is large enough to completely contain the unit circle. Project each point onto the unit circle. The distribution is not uniform since more points fall on a line from the origin to a vertex of the square than fall on a line from the origin to the midpoint of an edge of the square due to the difference in length. To solve this problem, discard all points outside the unit circle and project the remaining points onto the circle.

One might generalize this technique in the obvious way to higher dimensions. However, the ratio of the volume of a d -dimensional unit sphere to the volume of a d -dimensional 2 by 2 cube decreases rapidly making the process impractical for high dimensions since almost no points will lie inside the sphere. The solution is to generate a point each of whose coordinates is a Gaussian variable. The probability distribution for a point $(x_1, x_2, \dots, x_d)$ is given by

$$p(x_1, x_2, \dots, x_d) = \frac{1}{(2\pi)^{\frac{d}{2}}} e^{-\frac{x_1^2 + x_2^2 + \dots + x_d^2}{2}}$$

and is spherically symmetric. Normalizing the vector $\mathbf{x} = (x_1, x_2, \dots, x_d)$ to a unit vector gives a distribution that is uniform over the sphere. Note that once the vector is normalized, its coordinates are no longer statistically independent.

2.6 Gaussians in High Dimension

A 1-dimensional Gaussian has its mass close to the origin. However, as the dimension is increased something different happens. The d -dimensional spherical Gaussian with zero mean and variance σ^2 has density function

$$p(\mathbf{x}) = \frac{1}{(2\pi)^{d/2} \sigma^d} \exp\left(-\frac{|\mathbf{x}|^2}{2\sigma^2}\right).$$

The value of the Gaussian is maximum at the origin, but there is very little volume there. When $\sigma^2 = 1$, integrating the probability density over a unit sphere centered at the origin yields nearly zero mass since the volume of a unit sphere is negligible. In fact, one needs to increase the radius of the sphere to $\sqrt{d}$ before there is a significant nonzero volume and hence a nonzero probability mass. If one increases the radius beyond $\sqrt{d}$, the integral ceases to increase, even though the volume increases, since the probability density is dropping off at a much higher rate. The natural scale for the Gaussian is in units of $\sigma\sqrt{d}$.

Expected squared distance of a point from the center of a Gaussian

Consider a d -dimensional Gaussian centered at the origin with variance σ^2 . For a point $\mathbf{x} = (x_1, x_2, \dots, x_d)$ chosen at random from the Gaussian, the expected squared length of $\mathbf{x}$ is

$$E(x_1^2 + x_2^2 + \dots + x_d^2) = d E(x_1^2) = d\sigma^2.$$

For large d , the value of the squared length of $\mathbf{x}$ is tightly concentrated about its mean and thus, although $E(x^2) \neq E^2(x)$, $E(x) \approx \sqrt{E(x^2)}$. We call the square root of the expected squared distance $\sigma\sqrt{d}$ the radius of the Gaussian. In the rest of this section, we consider spherical Gaussians with $\sigma = 1$. All results can be scaled up by σ .

The probability mass of a unit-variance Gaussian as a function of the distance from its center is given by $r^{d-1}e^{-r^2/2}$ times some constant normalization factor where r is the distance from the center and d is the dimension of the space. The probability mass function has its maximum at

$$r = \sqrt{d-1},$$

which can be seen from setting the derivative equal to zero.

$$\frac{\partial}{\partial r} r^{d-1} e^{-\frac{r^2}{2}} = (d-1)r^{d-2} e^{-\frac{r^2}{2}} - r^d e^{-\frac{r^2}{2}} = 0$$

Dividing by $r^{d-2}e^{-\frac{r^2}{2}}$, yields $r^2 = d-1$.

Width of the annulus

The Gaussian distribution in high dimensions, centered at the origin, has its maximum value at the origin. However, there is no probability mass in a sphere of radius one centered at the origin since the sphere has zero volume. In fact, there is no probability mass until one gets sufficiently far from the origin so a sphere of that radius has nonzero volume. This occurs at radius $\sqrt{d}$. Once one gets a little farther from the origin there is again no probability mass since the probability distribution is dropping exponentially fast and the volume of the sphere is only increasing polynomially fast. All the probability mass is in a narrow annulus of radius approximately $\sqrt{d}$. In Section 2.7 we prove that for any positive real number $\beta < \sqrt{d}$, all but $3e^{-c\beta^2}$ of the mass lies within the annulus $\sqrt{d} - \beta \leq r \leq \sqrt{d} + \beta$. See Theorem ??.

Separating Gaussians

Gaussians are often used to model data. A common stochastic model is the mixture model where one hypothesizes that the data is generated from a convex combination of simple probability densities. An example is two Gaussian densities $p_1(\mathbf{x})$ and $p_2(\mathbf{x})$ where data is drawn from the mixture $p(\mathbf{x}) = w_1 p_1(\mathbf{x}) + w_2 p_2(\mathbf{x})$ with positive weights w_1 and w_2 summing to one. Assume that p_1 and p_2 are spherical with unit variance. If their means are very close, then given data from the mixture, one cannot tell for each data point whether it came from p_1 or p_2 . The question arises as to how much separation is needed between the means to determine which Gaussian generated which data point. We will see that a separation of $\Omega(d^{1/4})$ suffices. The algorithm to separate two Gaussians is simple. Calculate the distance between all pairs of points. Points whose distance apart is smaller are from the same Gaussian, points whose distance is larger are from different

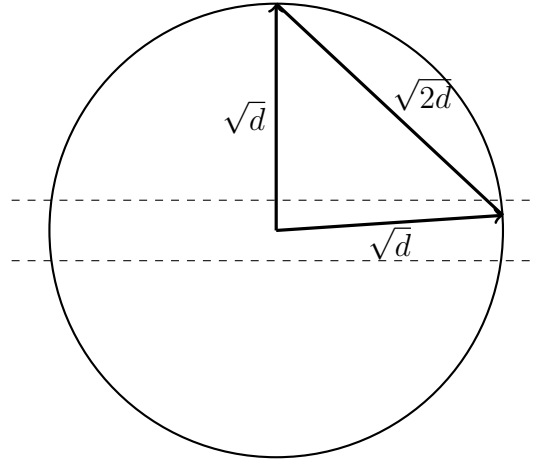


Figure 2.11: Two randomly chosen points in high dimension are almost surely nearly orthogonal.

Gaussians. Later, we will see that with more sophisticated algorithms, even a separation of $\Omega(1)$ suffices.

Consider two spherical unit-variance Gaussians. From Theorem ??, most of the probability mass of each Gaussian lies on an annulus of width $O(1)$ at radius $\sqrt{d-1}$. Also $e^{-|\mathbf{x}|^2/2} = \prod_i e^{-x_i^2/2}$ and almost all of the mass is within the slab $\{\mathbf{x} \mid -c \leq x_1 \leq c\}$, for $c \in O(1)$. Pick a point $\mathbf{x}$ from the first Gaussian. After picking $\mathbf{x}$, rotate the coordinate system to make the first axis point towards $\mathbf{x}$. Independently pick a second point $\mathbf{y}$ also from the first Gaussian. The fact that almost all of the mass of the Gaussian is within the slab $\{\mathbf{x} \mid -c \leq x_1 \leq c, c \in O(1)\}$ at the equator implies that $\mathbf{y}$'s component along $\mathbf{x}$'s direction is $O(1)$ with high probability. Thus, $\mathbf{y}$ is nearly perpendicular to $\mathbf{x}$. So, $|\mathbf{x} - \mathbf{y}| \approx \sqrt{|\mathbf{x}|^2 + |\mathbf{y}|^2}$. See Figure 2.11. More precisely, since the coordinate system has been rotated so that $\mathbf{x}$ is at the North Pole, $\mathbf{x} = (\sqrt{d} \pm O(1), 0, \dots, 0)$. Since $\mathbf{y}$ is almost on the equator, further rotate the coordinate system so that the component of $\mathbf{y}$ that is perpendicular to the axis of the North Pole is in the second coordinate. Then $\mathbf{y} = (O(1), \sqrt{d} \pm O(1), 0, \dots, 0)$. Thus,

$$(\mathbf{x} - \mathbf{y})^2 = d \pm O(\sqrt{d}) + d \pm O(\sqrt{d}) = 2d \pm O(\sqrt{d})$$

and $|\mathbf{x} - \mathbf{y}| = \sqrt{2d} \pm O(1)$.

Given two spherical unit variance Gaussians with centers $\mathbf{p}$ and $\mathbf{q}$ separated by a distance δ , the distance between a randomly chosen point $\mathbf{x}$ from the first Gaussian and a randomly chosen point $\mathbf{y}$ from the second is close to $\sqrt{\delta^2 + 2d}$, since $\mathbf{x} - \mathbf{p}$, $\mathbf{p} - \mathbf{q}$, and $\mathbf{q} - \mathbf{y}$ are nearly mutually perpendicular. Pick $\mathbf{x}$ and rotate the coordinate system so that $\mathbf{x}$ is at the North Pole. Let $\mathbf{z}$ be the North Pole of the sphere approximating the second Gaussian. Now pick $\mathbf{y}$. Most of the mass of the second Gaussian is within $O(1)$

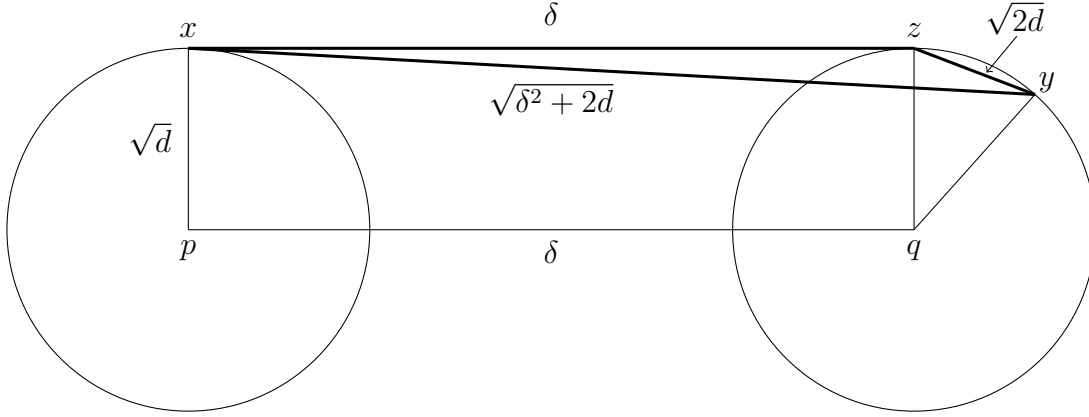


Figure 2.12: Distance between a pair of random points from two different unit spheres approximating the annuli of two Gaussians.

of the equator perpendicular to $\mathbf{q} - \mathbf{z}$. Also, most of the mass of each Gaussian is within distance $O(1)$ of the respective equators perpendicular to the line $\mathbf{q} - \mathbf{p}$. See Figure 2.12. Thus,

$$\begin{aligned} |\mathbf{x} - \mathbf{y}|^2 &\approx \delta^2 + |\mathbf{z} - \mathbf{q}|^2 + |\mathbf{q} - \mathbf{y}|^2 \\ &= \delta^2 + 2d \pm O(\sqrt{d}). \end{aligned}$$

To ensure that the distance between two points picked from the same Gaussian are closer to each other than two points picked from different Gaussians requires that the upper limit of the distance between a pair of points from the same Gaussian is at most the lower limit of distance between points from different Gaussians. This requires that $\sqrt{2d} + O(1) \leq \sqrt{2d + \delta^2} - O(1)$ or $2d + O(\sqrt{d}) \leq 2d + \delta^2$, which holds when $\delta \in \Omega(d^{1/4})$. Thus, mixtures of spherical Gaussians can be separated, provided their centers are separated by more than $d^{1/4}$. One can actually separate Gaussians where the centers are much closer. Chapter 4 contains an algorithm that separates a mixture of k spherical Gaussians whose centers are much closer.

Algorithm for separating points from two Gaussians

Calculate all pairwise distances between points. The cluster of smallest pairwise distances must come from a single Gaussian. Remove these points. The remaining points come from the second Gaussian.

Fitting a single spherical Gaussian to data

Given a set of sample points, $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$, in a d -dimensional space, we wish to find the spherical Gaussian that best fits the points. Let F be the unknown Gaussian with

mean $\boldsymbol{\mu}$ and variance σ^2 in each direction. The probability of picking these points when sampling according to F is given by

$$c \exp \left(- \frac{(\mathbf{x}_1 - \boldsymbol{\mu})^2 + (\mathbf{x}_2 - \boldsymbol{\mu})^2 + \cdots + (\mathbf{x}_n - \boldsymbol{\mu})^2}{2\sigma^2} \right)$$

where the normalizing constant c is the reciprocal of $\left[\int e^{-\frac{|\mathbf{x}-\boldsymbol{\mu}|^2}{2\sigma^2}} dx \right]^n$. In integrating from $-\infty$ to ∞ , one could shift the origin to $\boldsymbol{\mu}$ and thus c is $\left[\int e^{-\frac{|\mathbf{x}|^2}{2\sigma^2}} dx \right]^{-n} = \frac{1}{(2\pi)^{\frac{n}{2}}}$ and is independent of $\boldsymbol{\mu}$.

The *Maximum Likelihood Estimator* (MLE) of F , given the samples $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$, is the F that maximizes the above probability.

Lemma 2.8 *Let $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n\}$ be a set of n points in d -space. Then $(\mathbf{x}_1 - \boldsymbol{\mu})^2 + (\mathbf{x}_2 - \boldsymbol{\mu})^2 + \cdots + (\mathbf{x}_n - \boldsymbol{\mu})^2$ is minimized when $\boldsymbol{\mu}$ is the centroid of the points $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$, namely $\boldsymbol{\mu} = \frac{1}{n}(\mathbf{x}_1 + \mathbf{x}_2 + \cdots + \mathbf{x}_n)$.*

Proof: Setting the gradient of $(\mathbf{x}_1 - \boldsymbol{\mu})^2 + (\mathbf{x}_2 - \boldsymbol{\mu})^2 + \cdots + (\mathbf{x}_n - \boldsymbol{\mu})^2$ with respect $\boldsymbol{\mu}$ to zero yields

$$-2(\mathbf{x}_1 - \boldsymbol{\mu}) - 2(\mathbf{x}_2 - \boldsymbol{\mu}) - \cdots - 2(\mathbf{x}_n - \boldsymbol{\mu}) = 0.$$

Solving for $\boldsymbol{\mu}$ gives $\boldsymbol{\mu} = \frac{1}{n}(\mathbf{x}_1 + \mathbf{x}_2 + \cdots + \mathbf{x}_n)$. ■

To determine the maximum likelihood estimate of σ^2 for F , set $\boldsymbol{\mu}$ to the true centroid. Next, we show that σ is set to the standard deviation of the sample. Substitute $\nu = \frac{1}{2\sigma^2}$ and $a = (\mathbf{x}_1 - \boldsymbol{\mu})^2 + (\mathbf{x}_2 - \boldsymbol{\mu})^2 + \cdots + (\mathbf{x}_n - \boldsymbol{\mu})^2$ into the formula for the probability of picking the points $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$. This gives

$$\frac{e^{-a\nu}}{\left[\int_x e^{-x^2\nu} dx \right]^n}.$$

Now, a is fixed and ν is to be determined. Taking logs, the expression to maximize is

$$-a\nu - n \ln \left[\int_x e^{-\nu x^2} dx \right].$$

To find the maximum, differentiate with respect to ν , set the derivative to zero, and solve for σ . The derivative is

$$-a + n \frac{\int_x |x|^2 e^{-\nu x^2} dx}{\int_x e^{-\nu x^2} dx}.$$

Setting $y = |\sqrt{\nu}\mathbf{x}|$ in the derivative, yields

$$-a + \frac{n}{\nu} \frac{\int_y y^2 e^{-y^2} dy}{\int_y e^{-y^2} dy}.$$

Since the ratio of the two integrals is the expected distance squared of a d -dimensional spherical Gaussian of standard deviation $\frac{1}{\sqrt{2}}$ to its center, and this is known to be $\frac{d}{2}$, we get $-a + \frac{nd}{2\nu}$. Substituting σ^2 for $\frac{1}{2\nu}$ gives $-a + nd\sigma^2$. Setting $-a + nd\sigma^2 = 0$ shows that the maximum occurs when $\sigma = \frac{\sqrt{a}}{\sqrt{nd}}$. Note that this quantity is the square root of the average coordinate distance squared of the samples to their mean, which is the standard deviation of the sample. Thus, we get the following lemma.

Lemma 2.9 *The maximum likelihood spherical Gaussian for a set of samples is the one with center equal to the sample mean and standard deviation equal to the standard deviation of the sample from the true mean.*

Let $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ be a sample of points generated by a Gaussian probability distribution. $\boldsymbol{\mu} = \frac{1}{n}(\mathbf{x}_1 + \mathbf{x}_2 + \dots + \mathbf{x}_n)$ is an unbiased estimator of the expected value of the distribution. However, if in estimating the variance from the sample set, we use the estimate of the expected value rather than the true expected value, we will not get an unbiased estimate of the variance, since the sample mean is not independent of the sample set. One should use $\boldsymbol{\mu} = \frac{1}{n-1}(\mathbf{x}_1 + \mathbf{x}_2 + \dots + \mathbf{x}_n)$ when estimating the variance. See Section 11.4.8 of the appendix.

2.7 Bounds on Tail Probability

Markov's inequality bounds the tail probability of a nonnegative random variable x based only on its expectation. For $a > 0$,

$$\text{Prob}(x > a) \leq \frac{E(x)}{a}.$$

As a grows, the bound drops off as $1/a$. Given the second moment of x , Chebyshev's inequality, which does not assume x is a nonnegative random variable, gives a tail bound falling off as $1/a^2$

$$\text{Prob}(|x - E(x)| \geq a) \leq \frac{E\left((x - E(x))^2\right)}{a^2}.$$

Higher moments yield bounds by applying either of these two theorems. For example, if r is a nonnegative even integer, then x^r is a nonnegative random variable even if x takes on negative values. Applying Markov's inequality to x^r ,

$$\text{Prob}(|x| \geq a) = \text{Prob}(x^r \geq a^r) \leq \frac{E(x^r)}{a^r},$$

a bound that falls off as $1/a^r$. The larger the r , the greater the rate of fall, but a bound on $E(x^r)$ is needed to apply this technique.

For a random variable x that is the sum of a large number of independent random variables, $x_1, x_2, \dots, x_n$, one can derive bounds on $E(x^r)$ for high even r . There are many situations where the sum of a large number of independent random variables arises. For example, x_i may be the amount of a good that the i^{th} consumer buys, the length of the i^{th} message sent over a network, or the indicator random variable of whether the i^{th} record in a large database has a certain property. Each x_i is modeled by a simple probability distribution. Gaussian, exponential (probability density at any $t > 0$ is e^{-t}), or binomial distributions are typically used, in fact, respectively in the three examples here. If the x_i have 0-1 distributions, there are a number of theorems called Chernoff bounds, bounding the tails of $x = x_1 + x_2 + \dots + x_n$, typically proved by the so-called moment-generating function method (see Section 11.5 of the appendix). But exponential and Gaussian random variables are not bounded and these methods do not apply. However, good bounds on the moments of these two distributions are known. Indeed, for any integer $s > 0$, the s^{th} moment for the unit variance Gaussian and the exponential are both at most $s!$.

Given bounds on the moments of individual x_i the following theorem proves moment bounds on their sum. We use this theorem to derive tail bounds not only for sums of 0-1 random variables, but also Gaussians, exponentials, Poisson, etc.

The gold standard for tail bounds is the central limit theorem for independent, identically distributed random variables $x_1, x_2, \dots, x_n$ with zero mean and $\text{Var}(x_i) = \sigma^2$ that states as $n \rightarrow \infty$ the distribution of $x = (x_1 + x_2 + \dots + x_n)/\sqrt{n}$ tends to the Gaussian density with zero mean and variance σ^2 . Loosely, this says that in the limit, the tails of $x = (x_1 + x_2 + \dots + x_n)/\sqrt{n}$ are bounded by that of a Gaussian with variance σ^2 . But this theorem is only in the limit, whereas, we prove a bound that applies for all n .

In the following theorem, x is the sum of n independent, not necessarily identically distributed, random variables $x_1, x_2, \dots, x_n$, each of zero mean and variance at most σ^2 . By the central limit theorem, in the limit the probability density of x goes to that of the Gaussian with variance at most $n\sigma^2$. In a limit sense, this implies an upper bound of $ce^{-a^2/(2n\sigma^2)}$ for the tail probability $\text{Prob}(|x| > a)$ for some constant c . The following theorem assumes bounds on higher moments, but asserts a quantitative upper bound of $3e^{-a^2/(8n\sigma^2)}$ on the tail probability, not just in the limit, but for every n . We will apply this theorem to get tail bounds on sums of Gaussian, binomial, and power law distributed random variables.

Theorem 2.10 *Let $x = x_1 + x_2 + \dots + x_n$, where $x_1, x_2, \dots, x_n$ are mutually independent random variables with zero mean and variance at most σ^2 . If for $3 \leq s \leq (a^2/4n\sigma^2)$, $|E(x_i^s)| \leq \sigma^2 s!$, then for $0 \leq a \leq \sqrt{2n\sigma^2}$,*

$$\text{Prob}(|x_1 + x_2 + \dots + x_n| \geq a) \leq 3e^{-a^2/(8n\sigma^2)}.$$

Proof: We first prove an upper bound on $E(x^r)$ for any even positive integer r and then use Markov's inequality as discussed earlier. Expand $(x_1 + x_2 + \cdots + x_n)^r$.

$$\begin{aligned} (x_1 + x_2 + \cdots + x_n)^r &= \sum \binom{r}{r_1, r_2, \dots, r_n} x_1^{r_1} x_2^{r_2} \cdots x_n^{r_n} \\ &= \sum \frac{r!}{r_1! r_2! \cdots r_n!} x_1^{r_1} x_2^{r_2} \cdots x_n^{r_n} \end{aligned}$$

where the r_i range over all nonnegative integers summing to r . By independence

$$E(x^r) = \sum \frac{r!}{r_1! r_2! \cdots r_n!} E(x_1^{r_1}) E(x_2^{r_2}) \cdots E(x_n^{r_n}).$$

If in a term, any $r_i = 1$, the term is zero since $E(x_i) = 0$. Assume henceforth that $(r_1, r_2, \dots, r_n)$ runs over sets of nonzero r_i summing to r where each nonzero r_i is at least two. There are at most $r/2$ nonzero r_i in each set. Since $|E(x_i^{r_i})| \leq \sigma^2 r_i!$,

$$E(x^r) \leq r! \sum_{(r_1, r_2, \dots, r_n)} \sigma^{2(\text{number of nonzero } r_i \text{ in set})}.$$

Collect terms of the summation with t nonzero r_i for $t = 1, 2, \dots, r/2$. There are $\binom{n}{t}$ subsets of $\{1, 2, \dots, n\}$ of cardinality t . Once a subset is fixed as the set of t values of i with nonzero r_i , set each of the $r_i \geq 2$. That is, allocate two to each of the r_i and then allocate the remaining $r - 2t$ to the t r_i arbitrarily. The number of such allocations is just $\binom{r-2t+t-1}{t-1} = \binom{r-t-1}{t-1}$. So,

$$E(x^r) \leq r! \sum_{t=1}^{r/2} f(t), \quad \text{where} \quad f(t) = \binom{n}{t} \binom{r-t-1}{t-1} \sigma^{2t}.$$

Thus $f(t) \leq h(t)$, where $h(t) = \frac{(n\sigma^2)^t}{t!} 2^{r-t-1}$. In the hypotheses of the theorem $a \leq \sqrt{2} n\sigma^2$ and $s \leq \frac{a^2}{4n\sigma^2}$. Thus r is at most $n\sigma^2/2$. For $t \leq r/2$, increasing t by one, increases $h(t)$ by at least $n\sigma^2/(2t)$, which is at least two. This gives

$$E(x^r) = r! \sum_{t=1}^{r/2} f(t) \leq r! h(r/2) (1 + \frac{1}{2} + \frac{1}{4} + \cdots) \leq \frac{r!}{(r/2)!} 2^{r/2} (n\sigma^2)^{r/2}.$$

Applying Markov inequality,

$$\text{Prob}(|x| > a) = \text{Prob}(|x|^r > a^r) \leq \frac{r! (n\sigma^2)^{r/2} 2^{r/2}}{(r/2)! a^r} = g(r).$$

For even r , $g(r)/g(r-2) = \frac{4(r-1)n\sigma^2}{a^2}$ and so $g(r)$ decreases as long as $r-1 \leq a^2/(4n\sigma^2)$. Taking r to be the largest even integer less than or equal to $a^2/(4n\sigma^2)$, the tail probability is at most $e^{-r/2}$, which is at most $e \cdot e^{-a^2/(8n\sigma^2)} \leq 3 \cdot e^{-a^2/(8n\sigma^2)}$, proving the theorem. ■

2.8 Applications of the tail bound

Calculation of width of the Gaussian annulus

Let $(y_1, y_2, \dots, y_d)$ be a unit variance Gaussian centered at the origin. We argue that the mass of the Gaussian is in a narrow annulus of radius approximately $\sqrt{d}$. It is easier to deal with squared distance to the origin rather than distance. Thus, we ask what is the probability that $|y_1^2 + y_2^2 + \dots + y_d^2 - d| \geq \beta$? Let $x_i = y_i^2 - 1$ and change the question to what is the probability that $|x_1 + x_2 + \dots + x_d| \geq \beta$ to which we can apply Theorem 2.10.

Theorem 2.10 requires bounds on the moments of the x_i . For $|y_i| \leq 1$, $|x_i|^s \leq 1$ and for $|y_i| \geq 1$, $|x_i|^s \leq |y_i|^{2s}$. Thus

$$\begin{aligned} |E(x_i^s)| &= E(|x_i|^s) \leq E(1 + y_i^{2s}) = 1 + E(y_i^{2s}) \\ &= 1 + \sqrt{\frac{2}{\pi}} \int_0^\infty y^{2s} e^{-y^2/2} dy \end{aligned}$$

Using the substitution $y^2 = 2z$,

$$\begin{aligned} |E(x_i^s)| &= 1 + \frac{2^s}{\sqrt{\pi}} \int_0^\infty z^s z^{s-(1/2)} e^{-z} dz \\ &\leq 2^s s!. \end{aligned}$$

The last inequality is from the Gamma integral.

Theorem 2.10 requires $|E(x_i^s)| \leq \sigma^2(z_i)s!$ not $2^s s!$. Let $z_i = \frac{x_i}{2}$ and apply the theorem to

$$|z_1 + z_2 + \dots + z_d| \geq \frac{c}{2}.$$

From the above $|E(z_i^s)| = |E(\frac{x_i^s}{2^s})| \leq s!$. and $\sigma^2(z_i) \leq 2$. Since $|E(z_i^s)| \leq s!$, the hypothesis of Theorem 2.10 is satisfied.

Theorem 2.11 *For a d -dimensional unit variance spherical Gaussian, for any positive real number $c \leq \sqrt{d}$, all but $3e^{-\frac{c^2}{64}}$ of the mass lies within the annulus $\sqrt{d}-c \leq r \leq \sqrt{d}+c$.*

Proof: Let r be the distance to a point generated by the Gaussian. If $|r - \sqrt{d}| \geq c$, then since $|r + \sqrt{d}| \geq \sqrt{d}$,

$$|r^2 - d| = |r - \sqrt{d}||r + \sqrt{d}| \geq c\sqrt{d}.$$

Thus

$$|y_1^2 + y_2^2 + \dots + y_d^2 - d| \geq c\sqrt{d}$$

and hence

$$|x_1 + x_2 + \dots + x_d| \geq c\sqrt{d}$$

or

$$|z_1 + z_2 + \dots + z_d| \geq \frac{c\sqrt{d}}{2}.$$

Applying Theorem 2.10 where $\sigma^2 = 2$ and $n = d$, this occurs with probability less than or equal to $3e^{-\frac{c^2}{64}}$. ■

Chernoff Bounds

Chernoff bounds deal with sums of Bernoulli random variables. Here we apply Theorem 2.10 to derive similar bounds.

Theorem 2.12 *Suppose $y_1, y_2, \dots, y_n$ are independent 0-1 random variables with $E(y_i) = p$ for all i . Let $y = y_1 + y_2 + \dots + y_n$. Then for any $c \in [0, 1]$,*

$$\text{Prob}(|y - E(y)| \geq cnp) \leq 3e^{-npc^2/8}.$$

Proof: Let $x_i = y_i - p$. Then, $E(x_i) = 0$ and $E(x_i^2) = E(y - p)^2 = p$. For $s \geq 3$,

$$\begin{aligned} |E(x_i^s)| &= |E(y_i - p)^s| \\ &= |p(1-p)^s + (1-p)(0-p)^s| \\ &= |p(1-p)((1-p)^{s-1} + (-p)^{s-1})| \\ &\leq p. \end{aligned}$$

Apply Theorem 2.10 with $a = cnp$. Noting that $a < \sqrt{2} np$, completes the proof. ■

The appendix contains a different proof that uses a standard method based on moment-generating functions, which gives a better constant in the exponent.

Power Law Distributions

The power law distribution of order k where k is a positive integer is

$$f(x) = \frac{k-1}{x^k} \quad \text{for } x \geq 1.$$

If a random variable x has this distribution for $k \geq 4$, then

$$\mu = E(x) = \frac{k-1}{k-2} \quad \text{and} \quad \text{Var}(x) = \frac{k-1}{(k-2)^2(k-3)}.$$

Theorem 2.13 *Suppose y obeys a power law of order $k \geq 4$ and $x_1, x_2, \dots, x_n$ are independent random variables, each with the same distribution as $y - E(y)$. Let $x = x_1 + x_2 + \dots + x_n$. For any nonnegative $a \leq \frac{1}{10} \sqrt{\frac{n}{k}}$,*

$$\text{Prob}(|x| \geq a) \leq e^{-\frac{a^2}{8\text{var}(x)}}.$$

Proof: For integer s , the s^{th} moment of x_i , namely, $E(x_i^s)$, exists if and only if $s \leq k - 2$. For $s \leq k - 2$,

$$E(x_i^s) = (k - 1) \int_1^\infty \frac{(y - u)^s}{y^k} dy$$

Using the substitution of variable $z = \mu/y$

$$\frac{(y - u)^s}{y^k} = y^{s-k}(1 - z)^s = \frac{z^{k-s}}{u^{k-s}}(1 - z)^s$$

As y goes from 1 to ∞ , z goes from μ to 0, and $dz = -\frac{\mu}{y^2} dy$. Thus

$$\begin{aligned} E(x_i^s) &= (k - 1) \int_1^\infty \frac{(y - \mu)^s}{y^k} dy \\ &= \frac{k - 1}{\mu^{k-s-1}} \int_0^1 (1 - z)^s z^{k-s-2} dz + \frac{k - 1}{\mu^{k-s-1}} \int_1^\mu (1 - z)^s z^{k-s-2} dz. \end{aligned}$$

The first integral is just the standard integral of the beta function and its value is $\frac{s!(k-2-s)!}{(k-1)!}$. To bound the second integral, note that for $z \in [1, \mu]$, $|z - 1| \leq \frac{1}{k-2}$ and

$$z^{k-s-2} \leq \left(1 + (1/(k-2))\right)^{k-s-2} \leq e^{(k-s-2)/(k-2)} \leq e.$$

Apply Theorem 2.10 requires bounding $|E(x_i^s)|$ for $3 \leq s \leq \left\lfloor \frac{a^2}{4n\text{Var}(x_i)} \right\rfloor$. Since $a \leq \frac{1}{10} \sqrt{\frac{n}{k}}$, it follows that

$$\left\lfloor \frac{a^2}{4n\text{Var}(x_i)} \right\rfloor \leq \frac{n}{100k} \frac{1}{4n} \frac{(k-2)^2(k-3)}{k-1} \leq \frac{(k-2)^2(k-3)}{k(k-1)} \leq k-2.$$

So it suffices to prove that $|E(x_i^s)| \leq s!\text{Var}(x)$ for $3 \leq s \leq \dots, k-2$. If $k = 4$, s can go only up to 2 and there is nothing to prove. So assume $k \geq 5$. Since $\mu > 1$,

$$|E(x_i^s)| \leq \frac{(k-1)s!(k-2-s)!}{(k-1)!} + \frac{e(k-1)}{(k-2)^{s+1}} \leq s!\text{Var}(y) \left(\frac{1}{k-4} + \frac{e}{3!} \right) \leq s!\text{Var}(x).$$

Now, the theorem follows from Theorem 2.10. ■

2.9 Random Projection and Johnson-Lindenstrauss Theorem

Many high-dimensional problems, such as the nearest neighbor problem, can be sped up by projecting the data to a lower-dimensional subspace and solving the problem there. It would be convenient to have a projection to a lower-dimensional subspace that reduced all distances by the same common factor, thereby leaving the relative ordering of distances unchanged. The Johnson-Lindenstrauss theorem states that a projection to a random low-dimensional subspace has this property. In this section, we prove the Johnson-Lindenstrauss theorem and illustrate its application.

A random subspace of dimension one is a random line through the origin. A random subspace of dimension k is specified by picking a random line through the origin, then a second random line through the origin orthogonal to the first line, and then a third line orthogonal to the first two, etc. Their span is the random subspace.

Project a fixed unit-length vector $\mathbf{v}$ in d -dimensional space onto a random k -dimensional space. By the Pythagoras theorem, the length squared of a vector is the sum of the squares of its components. Intuitively, in a random direction the squared length of the vector should be about $\frac{1}{d}$ and so the squared length of the projection into a random k dimensional space should be about k/d . Thus, we would expect the length of the projection to be $\sqrt{\frac{k}{d}}$. The following theorem asserts that with high probability the length of the projection is very close to this quantity with failure probability exponentially small in k .

Theorem 2.14 (The Random Projection Theorem) *Let $\mathbf{v}$ be a fixed unit length vector in a d -dimensional space and let W be a random k -dimensional subspace. Let $\mathbf{w}$ be the projection of $\mathbf{v}$ onto W . For $0 \leq \varepsilon \leq 1$, $\text{Prob} \left(\left| |\mathbf{w}| - \sqrt{\frac{k}{d}} \right| \geq \varepsilon \sqrt{\frac{k}{d}} \right) \leq 3e^{-\frac{k\varepsilon^2}{64}}$.*

Proof: It is difficult to work with a random subspace. However, projecting a fixed vector onto a random subspace is the same as projecting a random vector onto a fixed subspace since one can rotate the coordinate system so that a set of basis vectors for the random subspace are the first k coordinate axes. The fixed vector then becomes a random vector. Thus, the probability distribution of $\mathbf{w}$ in the theorem is the same as the probability distribution of the vector obtained by taking a random unit length vector $\mathbf{z}$ and projecting $\mathbf{z}$ onto the fixed subspace U spanned by the first k coordinate vectors.

Pick a random vector $\mathbf{z}$ of length one by picking independent Gaussian random variables $x_1, x_2, \dots, x_d$, each with mean zero and variance one. Let $\mathbf{x} = (x_1, x_2, \dots, x_d)$ and $\mathbf{z} = \mathbf{x}/|\mathbf{x}|$. The vector $\mathbf{z}$ is a random vector of length one.

Let $\tilde{\mathbf{z}}$ be the projection of $\mathbf{z}$ onto U . We will prove that $|\tilde{\mathbf{z}}| \approx \sqrt{\frac{k}{d}}$ with high probability. Let $a = \sqrt{x_1^2 + x_2^2 + \dots + x_k^2}$ be the length of the projection of $\mathbf{x}$ and let $b = \sqrt{x_1^2 + x_2^2 + \dots + x_d^2}$ be the length of $\mathbf{x}$. Then the length of the projection of $\mathbf{z} = \frac{\mathbf{x}}{|\mathbf{x}|}$ is $|\tilde{\mathbf{z}}| = \frac{a}{b}$.

If $k\varepsilon^2 < 64$, then $3e^{-\frac{k\varepsilon^2}{64}} > 3e^{-1} > 1$, and there is nothing to prove since the upper bound on the probability that the projection deviates significantly from $\sqrt{\frac{k}{d}}$ asserted in the theorem is greater than one. Assume that $k\varepsilon^2 \geq 64$ which implies $\varepsilon \geq \frac{8}{\sqrt{k}}$. Define $c = \varepsilon\sqrt{k}/4$.

Applying Theorem 2.11 twice, all of the following inequalities hold with probability at

least $1 - 3e^{-\frac{k\varepsilon^2}{128}}$.

$$\sqrt{k} - c \leq a \leq \sqrt{k} + c \quad (2.5)$$

$$\sqrt{d} - c \leq b \leq \sqrt{d} + c. \quad (2.6)$$

From (2.5) and (2.6),

$$\frac{\sqrt{k} - c}{\sqrt{d} + c} \leq \frac{a}{b} \leq \frac{\sqrt{k} + c}{\sqrt{d} - c}$$

Thus, the length, $|w| = \frac{a}{b}$, of the projection is bounded between $\frac{\sqrt{k}-c}{\sqrt{d}+c}$ and $\frac{\sqrt{k}+c}{\sqrt{d}-c}$. In a moment we will show that $\frac{\sqrt{k}-c}{\sqrt{d}+c} \geq (1-\varepsilon)\frac{k}{d}$ and $\frac{\sqrt{k}+c}{\sqrt{d}-c} \leq (1+\varepsilon)\frac{k}{d}$. Thus, the length of the projection is bounded by

$$(1-\varepsilon)\frac{k}{d} \leq |w| \leq (1+\varepsilon)\frac{k}{d}$$

or $||w| - \frac{k}{d}| \leq \varepsilon \frac{k}{d}$ with probability less than or equal to $3e^{-\frac{k\varepsilon^2}{64}}$ completing the proof.

To see that $\frac{\sqrt{k}-c}{\sqrt{d}+c} \geq (1-\varepsilon)\frac{k}{d}$, multiply out $\frac{\sqrt{k}-c}{\sqrt{d}+c} \geq (1-\varepsilon)\frac{\sqrt{k}}{\sqrt{d}}$, to get

$$\sqrt{k}\sqrt{d} - c\sqrt{d} \geq (1-\varepsilon)\sqrt{k}\sqrt{d} + c(1-\varepsilon)\sqrt{k}.$$

Collecting terms, this is equivalent to $\varepsilon\sqrt{k}\sqrt{d} \geq c\sqrt{d} + c\sqrt{k}$. Substituting $c = \varepsilon\sqrt{k}/4$ and $k \leq d$, establishes that the inequality holds.

Similarly, to see that $\frac{\sqrt{k}+c}{\sqrt{d}-c} \geq (1-\varepsilon)\frac{k}{d}$, multiply out $\frac{\sqrt{k}+c}{\sqrt{d}-c} \leq (1+\varepsilon)\frac{\sqrt{k}}{\sqrt{d}}$, to get

$$(1+\varepsilon)\sqrt{k}\sqrt{d} - c(1+\varepsilon)\sqrt{k} \geq \sqrt{k}\sqrt{d} + c\sqrt{d}.$$

Collecting terms, this is equivalent to $\sqrt{k}\sqrt{d} \geq c(1+\varepsilon)\sqrt{k} + c\sqrt{d}$. Substituting $c = \varepsilon\sqrt{k}/4$ and $k \leq d$, establishes that the inequality holds. $\blacksquare$

The random projection theorem establishes that the probability of the length of the projection of a single vector differing significantly from its expected value is exponentially small in k , the dimension of the target subspace. By a union bound, the probability that any of $O(n^2)$ pairwise differences among n vectors differs significantly from their expected values is small, provided $k\varepsilon^2$ is $\Omega(\ln n)$. Thus, the projection to a random subspace preserves all relative pairwise distances between points in a set of n points. This is the content of the Johnson-Lindenstrauss theorem.

Theorem 2.15 (Johnson-Lindenstrauss Theorem) *For any $0 < \varepsilon < 1$ and any integer n , let k satisfy $k\varepsilon^2 \geq 192 \ln n$. For any set P of n points in R^d , a random projection f mapping $f : R^d \rightarrow R^k$ has the property that for all $\mathbf{u}$ and $\mathbf{v}$ in P with probability at least $1 - (1/n)$,*

$$(1-\varepsilon)\sqrt{\frac{k}{d}}|\mathbf{u} - \mathbf{v}| \leq |f(\mathbf{u}) - f(\mathbf{v})| \leq (1+\varepsilon)\sqrt{\frac{k}{d}}|\mathbf{u} - \mathbf{v}|.$$

Proof: Applying the random projection theorem (Theorem 2.14), for any fixed $\mathbf{u}$ and $\mathbf{v}$, the probability that $|f(\mathbf{u}) - f(\mathbf{v})|$ is outside the range

$$\left[(1 - \varepsilon) \sqrt{\frac{k}{d}} |\mathbf{u} - \mathbf{v}|, (1 + \varepsilon) \sqrt{\frac{k}{d}} |\mathbf{u} - \mathbf{v}| \right]$$

is at most

$$3e^{-\frac{k\varepsilon^2}{64}} \leq \frac{3}{n^3}$$

for $k \geq \frac{3 \times 64 \ln n}{\varepsilon^2}$. By the union bound, the probability that some pair has a large distortion is less than $\binom{n}{2} \times \frac{3}{n^3} \leq \frac{1}{n}$. ■

Remark: It is important to note that the conclusion of Theorem 2.15 asserts for all $\mathbf{u}$ and $\mathbf{v}$ in P , not just for most $\mathbf{u}$ and $\mathbf{v}$. The weaker assertion for most $\mathbf{u}$ and $\mathbf{v}$ is not that useful, since we do not know which $\mathbf{v}$ might end up being the closest point to $\mathbf{u}$ and an assertion for most may not cover the particular $\mathbf{v}$. A remarkable aspect of the theorem is that the number of dimensions in the projection is only dependent logarithmically on n . Since k is often much less than d , this is called a dimension reduction technique.

For the nearest neighbor problem, if the database has n_1 points and n_2 queries are expected during the lifetime, take $n = n_1 + n_2$ and project the database to a random k -dimensional space, where $k \geq \frac{192 \ln n}{\varepsilon^2}$. On receiving a query, project the query to the same subspace and compute nearby database points. The Johnson Lindenstrauss theorem says that with high probability this will yield the right answer whatever the query. Note that the exponentially small in k probability in Theorem 2.14 was useful here in making k only dependent on $\ln n$, rather than n .

2.10 Bibliographic Notes

The word vector model was introduced by Salton [SWY75]. Taylor series remainder material can be found in Whittaker and Watson 1990, pp. 95-96 and Section 11.8.4 of the appendix. There is vast literature on the Gaussian distribution, its properties, drawing samples according to it, etc. The reader can choose the level and depth according to his/her background. For Chernoff bounds and their applications, see [MU05] or [MR95b]. The proof here and the application to heavy-tailed distributions is simplified from [Kan09]. The original proof of the random projection theorem by Johnson and Lindenstrauss was complicated. Several authors used Gaussians to simplify the proof. See [Vem04] for details and applications of the theorem. The proof here is due to Das Gupta and Gupta [DG99].

2.11 Exercises

Exercise 2.1

1. Let x and y be independent random variables with uniform distribution in $[0, 1]$. What is the expected value $E(x)$, $E(x^2)$, $E(x - y)$, $E(xy)$, and $E((x - y)^2)$?
2. Let x and y be independent random variables with uniform distribution in $[-\frac{1}{2}, \frac{1}{2}]$. What is the expected value $E(x)$, $E(x^2)$, $E(x - y)$, $E(xy)$, and $E((x - y)^2)$?
3. What is the expected squared distance between two points generated at random inside a unit d -dimensional cube centered at the origin?
4. Randomly generate a number of points inside a d -dimensional unit cube centered at the origin and plot distance between and the angle between the vectors from the origin to the points for all pairs of points.

Exercise 2.2 Consider two random 0-1 vectors in high dimension. What is the angle between them?

Exercise 2.3 In Section 2.1 on properties of high-dimensional space, we state that the distance of a point to the center of a sphere in d -dimensions is likely to be between $1 - \frac{c}{\sqrt{d}}$ and 1. We also claim that the first coordinate of such a point is likely to be between $-\frac{c}{\sqrt{d}}$ and $\frac{c}{\sqrt{d}}$. Justify the role of d in these statements.

Exercise 2.4 Show that Markov's inequality is tight by showing the following:

1. For each of $a = 2, 3$, and 4 give a probability distribution for a nonnegative random variable x where $\text{Prob}(x \geq aE(x)) = \frac{1}{a}$.
2. For arbitrary $a \geq 1$ give a probability distribution for a nonnegative random variable x where $\text{Prob}(x \geq aE(x)) = \frac{1}{a}$.

Exercise 2.5 In what sense is Chebyshev's inequality tight?

Exercise 2.6 Consider the probability function $p(x) = c\frac{1}{x^4}$, $x \geq 1$, and generate 100 random samples. How close is the average of the samples to the expected value of x ?

Exercise 2.7 What is the formula for the incremental unit of area when using polar coordinates to integrate the surface area of a unit radius 3-dimensional sphere lying inside a circular cone, whose vertex is at the center of the sphere? What is the formula for the integral? What is the value of the integral, if the cone is 36° ?

Exercise 2.8 For what value of d does the volume, $V(d)$, of a d -dimensional unit sphere take on its maximum?

Hint: Consider the ratio $\frac{V(d)}{V(d-1)}$.

Exercise 2.9 Write a recurrence relation for $V(d)$ in terms of $V(d-1)$ by integrating using an incremental unit that is a disk of thickness dr .

Exercise 2.10 How does the volume of a sphere of radius two behave as the dimension of the space increases? What if the radius was larger than two but a constant independent of d ? What function of d would the radius need to be for a sphere of radius r to have approximately constant volume as the dimension increases?

Exercise 2.11 A 3-dimensional cube has vertices, edges, and faces. In a d -dimensional cube, these components are called faces. A vertex is a 0-dimensional face, an edge a 1-dimensional face, etc. For $0 \leq i \leq d$, how many i -dimensional faces does a d -dimensional hyper cube have?

Exercise 2.12 Consider a unit radius, circular cylinder in 3-dimensions of height one. The top of the cylinder could be an horizontal plane or we could have half of a circular sphere. Consider these two possibilities for a unit radius, circular cylinder in 4-dimensions. In each of the two cases, what is the surface area of the top face of the cylinder? You can use $V(d)$ for the volume of a unit radius, d -dimension sphere and $A(d)$ for the surface area of a unit radius, d -dimensional sphere. An infinite length, unit radius, circular cylinder in 4-dimensions would be the set $\{(x_1, x_2, x_3, x_4) | x_2^2 + x_3^2 + x_4^2 \leq 1\}$ where the coordinate x_1 is the axis.

Exercise 2.13 Consider vertices of a d -dimensional cube of width two centered at the origin. Vertices are the points $(\pm 1, \pm 1, \dots, \pm 1)$. Place a unit-radius sphere at each vertex. Each sphere fits in a cube of width two and thus no two spheres intersect. Show that the probability that a point of the cube picked at random will fall into one of the 2^d unit-radius spheres, centered at the vertices of the cube, goes to 0 as d tends to infinity.

Exercise 2.14 Place two unit-radius spheres in d -dimensions, one at $(-2, 0, 0, \dots, 0)$ and the other at $(2, 0, 0, \dots, 0)$. Give an upper bound on the probability that a random line through the origin will intersect the spheres.

Exercise 2.15 Let $\mathbf{x}$ be a random sample from the unit sphere $\{\mathbf{x} | |\mathbf{x}| \leq 1\}$ in d -dimensions with the origin as center.

1. What is the mean of the random variable $\mathbf{x}$? The mean, denoted $E(\mathbf{x})$, is the vector, whose i^{th} component is $E(x_i)$.
2. What is the component-wise variance of $\mathbf{x}$?
3. For any unit length vector $\mathbf{u}$, the variance of the real-valued random variable $\mathbf{u}^T \mathbf{x}$ is $\sum_{i=1}^d u_i^2 E(x_i^2)$. Note that the x_i are not independent. Using (2), simplify this expression for the variance of $\mathbf{x}$.

4. * Given two spheres in d -space, both of radius one whose centers are distance a apart, show that the volume of their intersection is at most

$$\frac{4e^{-\frac{a^2(d-1)}{8}}}{a\sqrt{d-1}}$$

times the volume of each sphere. Hint: Relate the volume of the intersection to the volume of a cap; then, use Lemma 2.6.

5. From (4), conclude that if the inter-center separation of the two spheres of radius r is $\Omega(r/\sqrt{d})$, then they share very small mass. Theoretically, at this separation, given randomly generated points from the two distributions, one inside each sphere, it is possible to tell which sphere contains which point, i.e., classify them into two clusters so that each cluster is exactly the set of points generated from one sphere. The actual classification requires an efficient algorithm to achieve this. Note that the inter-center separation required here goes to zero as d gets larger, provided the radius of the spheres remains the same. So, it is easier to tell apart spheres (of the same radii) in higher dimensions.
6. * In this part, you will carry out the same exercise for Gaussians. First, restate the shared mass of two spheres as $\int_{\mathbf{x} \in \text{space}} \min(f(x), g(x)) dx$, where f and g are just the uniform densities in the two spheres respectively. Make a similar definition for the shared mass of two spherical Gaussians. Using this, show that for two spherical Gaussians, each with standard deviation σ in every direction and with centers at distance a apart, the shared mass is at most $(c_1/a) \exp(-c_2 a^2 / \sigma^2)$, where c_1 and c_2 are constants. This translates to “if two spherical Gaussians have centers which are $\Omega(\sigma)$ apart, then they share very little mass”. Explain.

Exercise 2.16 Prove that $1 + x \leq e^x$ for all real x . For what values of x is the approximation $1 + x \approx e^x$ good?

Exercise 2.17 Derive an upper bound on $\int_{x=a}^{\infty} e^{-\frac{x^2}{2}} dx$ where a is a positive real. Discuss for what values of a this is a good bound.

Hint: Use $e^{-\frac{x^2}{2}} \leq \frac{x}{a} e^{-\frac{x^2}{2}}$ for $x \geq a$.

Exercise 2.18 Verify the formula $V(d) = 2 \int_0^1 V(d-1)(1-x_1^2)^{\frac{d-1}{2}} dx_1$ for $d = 1$ and $d = 2$ by integrating and comparing with $V(2) = \pi$ and $V(3) = \frac{4}{3}\pi$

Exercise 2.19 What is the volume of a radius r cylinder of height h in d -dimensions?

Exercise 2.20 Consider the upper hemisphere of a unit-radius sphere in d -dimensions. What is the height of the maximum volume cylinder that can be placed entirely inside the hemisphere? As you increase the height of the cylinder, you need to reduce the cylinder's radius so that it will lie entirely within the hemisphere.

Exercise 2.21 What is the volume of the maximum size d -dimensional hypercube that can be placed entirely inside a unit radius d -dimensional sphere?

Exercise 2.22 In showing that the volume of a unit sphere was near the equator we obtained an upper bound on the volume of the upper hemisphere above the slice of

$$\frac{1}{\epsilon(d-1)} e^{\frac{d-1}{2}\epsilon^2} V(d-1)$$

and a lower bound on the volume of the upper hemisphere of $\frac{1}{2\sqrt{d-1}}V(d-1)$. Show that for a radius r sphere these bounds become $\frac{r^{d+1}}{\epsilon(d-1)} e^{\frac{d-1}{2}(\frac{\epsilon}{r})^2} V(d-1)$ and $\frac{r^d}{2\sqrt{d-1}}V(d-1)$ and that the ratio is $\frac{2r}{\epsilon\sqrt{d-1}} e^{\frac{d-1}{2}(\frac{\epsilon}{r})^2}$.

Exercise 2.23 For a 1,000-dimensional unit-radius sphere centered at the origin, what fraction of the volume of the upper hemisphere is above the plane $x_1 = 0.1$? Above the plane $x_1 = 0.01$?

Exercise 2.24 Let $\{\mathbf{x} \mid |\mathbf{x}| \leq 1\}$ be a d -dimensional, unit radius sphere centered at the origin. What fraction of the volume is the set $\{(x_1, x_2, \dots, x_d) \mid \forall i \ |x_i| \leq \frac{1}{\sqrt{d}}\}$?

Exercise 2.25 Almost all of the volume of a sphere in high dimensions lies in a narrow slice of the sphere at the equator. However, the narrow slice is determined by the point on the surface of the sphere that is designated the North Pole. Explain how this can be true if several different locations are selected for the North Pole.

Exercise 2.26 Explain how the volume of a sphere in high dimensions can simultaneously be in a narrow slice at the equator and also be concentrated in a narrow annulus at the surface of the sphere.

Exercise 2.27 Project the vertices of a high-dimensional cube onto a line from $(0, 0, \dots, 0)$ to $(1, 1, \dots, 1)$. Argue that the “density” of the number of projected points (per unit distance) varies roughly as a Gaussian with variance $O(1)$ with the mid-point of the line as center.

Exercise 2.28

1. A unit cube has vertices, edges, faces, etc. How many k -dimensional objects are in a d -dimensional cube?
2. What is the surface area of a unit cube in d -dimensions?
3. What is the surface area of the cube if the length of each side was 2?
4. Prove that the volume of a unit cube is close to its surface.

Exercise 2.29 Define the equator of a d -dimensional unit cube to be the hyperplane $\left\{ \mathbf{x} \mid \sum_{i=1}^d x_i = \frac{d}{2} \right\}$.

1. Are the vertices of a unit cube concentrated close to the equator?
2. Is the volume of a unit cube concentrated close to the equator?
3. Is the surface area of a unit cube concentrated close to the equator?

Exercise 2.30 How large must ε be for 99% of the volume of a d -dimensional unit-radius sphere to lie in the shell of ε -thickness at the surface of the sphere?

Exercise 2.31 Calculate the ratio of area above the plane $x_1 = \epsilon$ of a unit radius sphere in d -dimensions for $\epsilon = 0.01, 0.02, 0.03, 0.04, 0.05$ and for $d = 100$ and $d = 1,000$. Also calculate the ratio for $\epsilon = 0.001$ and $d = 1,000$.

Exercise 2.32 1. What is the maximum size rectangle that can be fitted in a unit variance Gaussian?

2. What rectangle best approximates a unit variance Gaussian if one measure goodness of fit by how small the symmetric difference of the Gaussian and rectangle is.

Exercise 2.33 Generate 500 points uniformly at random on the surface of a unit-radius sphere in 50 dimensions. Then randomly generate five additional points. For each of the five new points, calculate a narrow band at the equator, assuming the point was the North Pole. How many of the 500 points are in each band corresponding to one of the five equators? How many of the points are in all five bands?

Exercise 2.34 We have claimed that a randomly generated point on a sphere lies near the equator of the sphere, wherever we place the North Pole. Is the same claim true for a randomly generated point on a cube? To test this claim, randomly generate ten ± 1 valued vectors in 128 dimensions. Think of these ten vectors as ten choices for the North Pole. Then generate some additional ± 1 valued vectors. To how many of the original vectors is each of the new vectors close to being perpendicular; that is, how many of the equators is each new vectors close to?

Exercise 2.35 Consider two random vectors in a high-dimensional space. Assume the vectors have been normalized so that their lengths are one and thus the points lie on a unit sphere. Assume one of the vectors is the North pole. Prove that the ratio of the area of a cone, with axis at the North Pole of fixed angle say 45° to the area of a hemisphere, goes to zero as the dimension increases. Thus, the probability that the angle between two random vectors is at most 45° goes to zero. How does this relate to the result that most of the volume is near the equator?

Exercise 2.36 Consider a slice of a 100-dimensional sphere that lies between two parallel planes, each equidistant from the equator and perpendicular to the line from the North to the South Pole. What percentage of the distance from the center of the sphere to the poles must the planes be to contain 95% of the surface area?

Exercise 2.37 Place n points at random on a d -dimensional unit-radius sphere. Assume d is large. Pick a random vector and let it define two parallel hyperplanes on opposite sides of the origin that are equal distance from the origin. How far apart can the hyperplanes be moved and still have the probability that none of the n points lands between them be at least .99?

Exercise 2.38 Project the surface area of a d -dimensional sphere of radius $\sqrt{d}$ onto a line through the center. For large d , give an intuitive argument that the projected surface area should behave like a Gaussian.

Exercise 2.39 * Consider the simplex

$$S = \{\mathbf{x} \mid x_i \geq 0, 1 \leq i \leq d; \sum_{i=1}^d x_i \leq 1\}.$$

For a random point $\mathbf{x}$ picked with uniform density from S , find $E(x_1 + x_2 + \cdots + x_d)$. Find the centroid of S .

Exercise 2.40 How would you sample uniformly at random from the parallelepiped

$$P = \{\mathbf{x} \mid \mathbf{0} \leq A\mathbf{x} \leq \mathbf{1}\},$$

where A is a given nonsingular matrix? How about from the simplex

$$\{\mathbf{x} \mid 0 \leq (A\mathbf{x})_1 \leq (A\mathbf{x})_2 \leq \cdots \leq (A\mathbf{x})_d \leq 1\}?$$

Your algorithms must run in polynomial time.

Exercise 2.41 Let G be a d -dimensional spherical Gaussian with variance $\frac{1}{2}$ centered at the origin. Derive the expected squared distance to the origin.

Exercise 2.42

1. Write a computer program that generates n points uniformly distributed over the surface of a unit-radius d -dimensional sphere.
2. Generate 200 points on the surface of a sphere in 50 dimensions.
3. Create several random lines through the origin and project the points onto each line. Plot the distribution of points on each line.

4. What does your result from (3) say about the surface area of the sphere in relation to the lines, i.e., where is the surface area concentrated relative to each line?

Exercise 2.43 *If one generates points in d -dimensions with each coordinate a unit variance Gaussian, the points will approximately lie on the surface of a sphere of radius $\sqrt{d}$.*

1. What is the distribution when the points are projected onto a random line through the origin?
2. If one uses a Gaussian with variance four, where in d -space will the points lie?

Exercise 2.44 *Randomly generate a 100 points on the surface of a sphere in 3-dimensions and in 100-dimensions. Create a histogram of all distances between the pairs of points in both cases.*

Exercise 2.45 *We have claimed that in high dimensions, a unit variance Gaussian centered at the origin has essentially zero probability mass in a unit-radius sphere centered at the origin. Show that as the variance of the Gaussian goes down, more and more of its mass is contained in the unit-radius sphere. How small must the variance be for 0.99 of the mass of the Gaussian to be contained in the unit-radius sphere?*

Exercise 2.46 *Consider two unit-radius spheres in d -dimensions whose centers are distance δ apart where δ is a constant independent of d . Let $\mathbf{x}$ be a random point on the surface of the first sphere and $\mathbf{y}$ a random point on the surface of the second sphere. Prove that the probability that $|\mathbf{x} - \mathbf{y}|^2$ is more than $2 + \delta^2 + a$, falls off exponentially with a .*

Exercise 2.47 * *Pick a point $\mathbf{x}$ uniformly at random from the following set in d -space:*

$$K = \{\mathbf{x} | x_1^4 + x_2^4 + \cdots + x_d^4 \leq 1\}.$$

1. Show that the probability that $x_1^4 + x_2^4 + \cdots + x_d^4 \leq \frac{1}{2}$ is $\frac{1}{2^{d/4}}$.
2. Show that with high probability, $x_1^4 + x_2^4 + \cdots + x_d^4 \geq 1 - O(1/d)$.
3. Show that with high probability, $|x_1| \leq O(1/d^{1/4})$.

Exercise 2.48 *Suppose there is an object moving at constant velocity along a straight line. You receive the gps coordinates corrupted by Gaussian noise every minute. How do you estimate the current position?*

Exercise 2.49 *Generate ten values by a Gaussian probability distribution with zero mean and variance one. What is the center determined by averaging the points? What is the variance? In estimating the variance, use both the real center and the estimated center. When using the estimated center to estimate the variance, use both $n = 10$ and $n = 9$. How do the three estimates compare?*

Exercise 2.50 Let $x_1, x_2, \dots, x_n$ be independent samples of a random variable $\mathbf{x}$ with mean m and variance σ^2 . Let $m_s = \frac{1}{n} \sum_{i=1}^n x_i$ be the sample mean. Suppose one estimates the variance using the sample mean rather than the true mean, that is,

$$\sigma_s^2 = \frac{1}{n} \sum_{i=1}^n (x_i - m_s)^2$$

Prove that $E(\sigma_s^2) = \frac{n-1}{n} \sigma^2$ and thus one should have divided by $n-1$ rather than n .

Hint: First calculate the variance of the sample mean and show that $\text{var}(m_s) = \frac{1}{n} \text{var}(\mathbf{x})$. Then calculate $E(\sigma_s^2) = E[\frac{1}{n} \sum_{i=1}^n (x_i - m_s)^2]$ by replacing $x_i - m_s$ with $(x_i - m) - (m_s - m)$.

Exercise 2.51 Suppose you want to estimate the unknown center of a Gaussian in d -space which has variance one in each direction. Show that $O(\log d / \varepsilon^2)$ random samples from the Gaussian are sufficient to get an estimate $\tilde{\mu}$ of the true center μ , so that with probability at least $99/100$,

$$|\mu - \tilde{\mu}|_\infty \leq \varepsilon.$$

How many samples are sufficient to ensure that

$$|\mu - \tilde{\mu}| \leq \varepsilon?$$

Exercise 2.52 Use the probability distribution $\frac{1}{\sqrt{2\pi}3} e^{-\frac{(x-5)^2}{2 \times 9}}$ to generate ten points.

(a) From the ten points estimate μ . How close is the estimate of μ to the true mean of 5?

(b) Using the true mean of 5, estimate σ^2 by the formula $\sigma^2 = \frac{1}{10} \sum_{i=1}^{10} (x_i - 5)^2$. How close is the estimate of σ^2 to the true variance of 9?

(c) Using your estimate of the mean, estimate σ^2 by the formula $\sigma^2 = \frac{1}{10} \sum_{i=1}^{10} (x_i - \hat{\mu})^2$. How close is the estimate of σ^2 to the true variance of 9?

(d) Using your estimate of the mean, estimate σ^2 by the formula $\sigma^2 = \frac{1}{9} \sum_{i=1}^{10} (x_i - \hat{\mu})^2$. How close is the estimate of σ^2 to the true variance of 9?

Exercise 2.53 The Cauchy distribution in one dimension is $\text{Prob}(x) = \frac{1}{c+x^2}$. What would happen if one tried to extend the distribution to higher dimensions by the formula $\text{Prob}(r) = \frac{1}{1+r^2}$, where r is the distance from the origin? What happens when you try to determine a normalization constant c ?

Exercise 2.54 Consider the power law probability density

$$p(x) = \frac{c}{\max(1, x^2)} = \begin{cases} c & 0 \leq x \leq 1 \\ \frac{c}{x^2} & x > 1 \end{cases}$$

over the nonnegative real line.

1. Determine the constant c .
2. For a nonnegative random variable x with this density, does $E(x)$ exist? How about $E(x^2)$?

Exercise 2.55 Consider d -space and the following density over the positive orthant:

$$p(\mathbf{x}) = \frac{c}{\max(1, |\mathbf{x}|^a)}.$$

Show that $a > d$ is necessary for this to be a proper density function. Show that $a > d + 1$ is a necessary condition for a (vector-valued) random variable $\mathbf{x}$ with this density to have an expected value $E(|\mathbf{x}|)$. What condition do you need if we want $E(|\mathbf{x}|^2)$ to exist?

Exercise 2.56 Assume you can generate a value uniformly at random in the interval $[0, 1]$. How would you generate a value according to the probability distribution $p(x)$?

Exercise 2.57 Let x be a random variable with probability density $\frac{1}{2}$ for $0 \leq x \leq 2$ and zero elsewhere. Use Markov's inequality to bound the probability that $x > 2$. Then, make use of $\text{Prob}(|x| > a) = \text{Prob}(x^2 > a^2)$ to get a tighter bound.

Exercise 2.58 Consider the probability distribution $p(x = 0) = 1 - \frac{1}{a}$ and $p(x = a) = \frac{1}{a}$. Plot the probability that x is greater than or equal to b as a function of b for the bound given by Markov's inequality and by Markov's inequality applied to x^2 and x^4 .

Exercise 2.59 Suppose $\mathbf{x}$ and $\mathbf{y}$ are two random 0-1 d -vectors. Show that with high probability the cosine of the angle between them is close to $\frac{1}{2}$. Hint: Model your proof after that of the random projection theorem.

Exercise 2.60 Generate 20 points uniformly at random on a 1,000-dimensional sphere of radius 100. Calculate the distance between each pair of points. Then, project the data onto subspaces of dimension $k=100, 50, 10, 5, 4, 3, 2, 1$ and calculate the sum of squared error between $\frac{k}{d}$ times the original distances and the new pair-wise distances for each of the above values of k .

Exercise 2.61 You are given two sets, P and Q , of n points each in n -dimensional space. Your task is to find the closest pair of points, one each from P and Q , i.e., find $\mathbf{x}$ in P and $\mathbf{y}$ in Q such that $|\mathbf{x} - \mathbf{y}|$ is minimum.

1. Show that this can be done in time $O(n^3)$.

2. Show how to do this with relative error 0.1% in time $O(n^2 \ln n)$, i.e., you must find a pair $\mathbf{x} \in P, \mathbf{y} \in Q$ so that the distance between them is, at most, 1.001 times the minimum possible distance. If the minimum distance is 0, you must find $\mathbf{x} = \mathbf{y}$.

Exercise 2.62 Given n data points in d -space, find a subset of k data points whose vector sum has the smallest length. You can try all $\binom{n}{k}$ subsets, compute each vector sum in time $O(kd)$ for a total time of $O(\binom{n}{k}kd)$. Show that we can replace d in the expression above by $O(k \ln n)$, if we settle for an answer with relative error .02%.

Exercise 2.63 To preserve pairwise distances between n data points in d space, we projected to a random $O(\ln n / \varepsilon^2)$ dimensional space. To save time in carrying out the projection, we may try to project to a space spanned by sparse vectors, vectors with only a few nonzero entries. that is, choose say $O(\ln n / \varepsilon^2)$ vectors at random, each with 100 nonzero components and project to the space spanned by them. Will this work (to preserve approximately all pairwise distances) ? Why?

Exercise 2.64 Create a list of the five most important things that you learned about high dimensions.

Exercise 2.65 Write a short essay whose purpose is to excite a college freshman to learn about high dimensions.

3 Random Graphs

Large graphs appear in many contexts such as the World Wide Web, the internet, social networks, journal citations, and other places. What is different about the modern study of large graphs from traditional graph theory and graph algorithms is that here one seeks statistical properties of these very large graphs rather than an exact answer to questions. This is akin to the switch physics made in the late 19th century in going from mechanics to statistical mechanics. Just as the physicists did, one formulates abstract models of graphs that are not completely realistic in every situation, but admit a nice mathematical development that can guide what happens in practical situations. Perhaps the most basic such model is the $G(n, p)$ model of a random graph. In this chapter, we study properties of the $G(n, p)$ model as well as other models.

3.1 The $G(n, p)$ Model

The $G(n, p)$ model, due to Erdős and Rényi, has two parameters, n and p . Here n is the number of vertices of the graph and p is the edge probability. For each pair of distinct vertices, v and w , p is the probability that the edge (v, w) is present. The presence of each edge is statistically independent of all other edges. The graph-valued random variable with these parameters is denoted by $G(n, p)$. When we refer to “the graph $G(n, p)$ ”, we mean one realization of the random variable. In many cases, p will be a function of n such as $p = d/n$ for some constant d . In this case, the expected degree of a vertex of the graph is $\frac{d}{n}(n-1) \approx d$. The interesting thing about the $G(n, p)$ model is that even though edges are chosen independently with no “collusion”, certain global properties of the graph emerge from the independent choices. For small p , with $p = d/n$, $d < 1$, each connected component in the graph is small. For $d > 1$, there is a giant component consisting of a constant fraction of the vertices. In addition, there is a rapid transition at the threshold $d = 1$. Below the threshold, the probability of a giant component is very small, and above the threshold, the probability is almost one.

The phase transition at the threshold $d = 1$ from very small $o(n)$ size components to a giant $\Omega(n)$ sized component is illustrated by the following example. Suppose the vertices represent people and an edge means the two people it connects know each other. Given a chain of connections, such as A knows B, B knows C, C knows D, ..., and Y knows Z, we say that A indirectly knows Z. Thus, all people belonging to a connected component of the graph indirectly know each other. Suppose each pair of people, independent of other pairs, tosses a coin that comes up heads with probability $p = d/n$. If it is heads, they know each other; if it comes up tails, they don't. The value of d can be interpreted as the expected number of people a single person directly knows. The question arises as to how large are sets of people who indirectly know each other?

If the expected number of people each person knows is more than one, then a giant component of people, all of whom indirectly know each other, will be present consisting

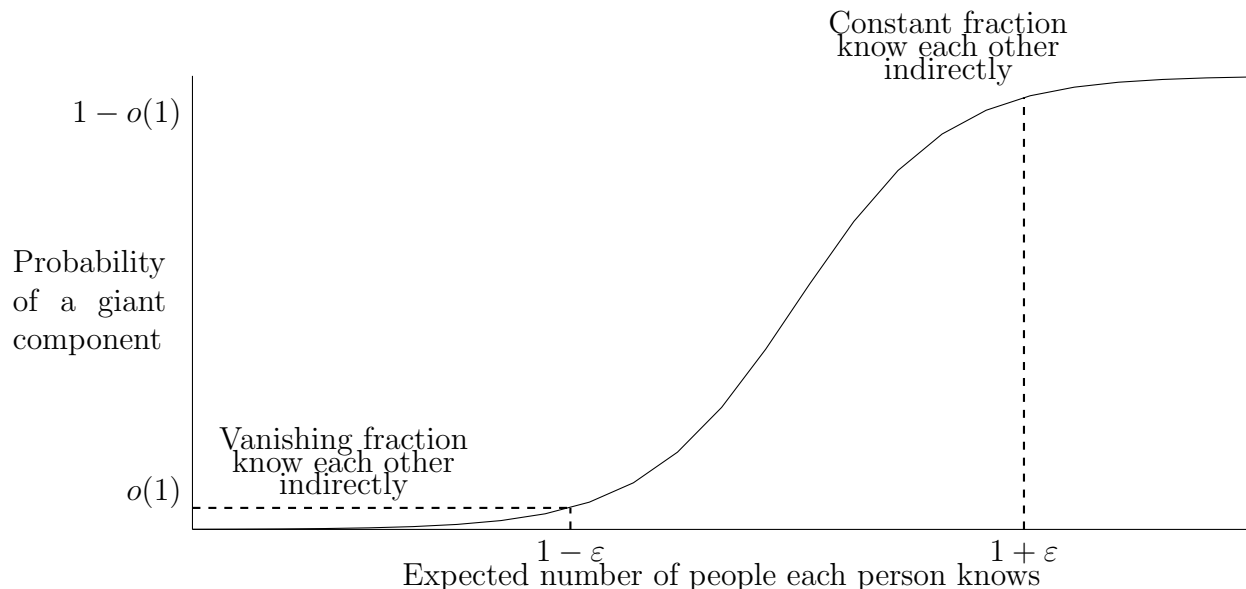


Figure 3.1: Probability of a giant component as a function of the expected number of people each person knows directly.

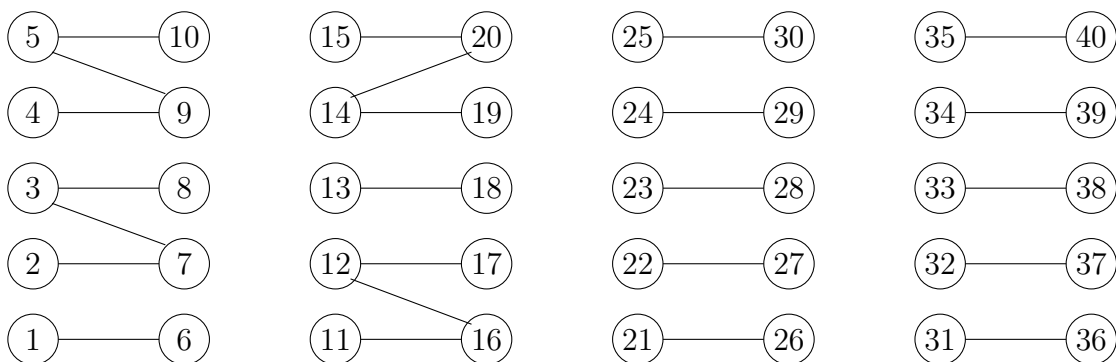
of a constant fraction of all the people. On the other hand, if in expectation, each person knows less than one person, the largest set of people who know each other indirectly is a vanishingly small fraction of the whole. Furthermore, the transition from the vanishing fraction to a constant fraction of the whole, happens abruptly between d slightly less than one to d slightly more than one. See Figure 3.1. Note that there is no global coordination of who knows whom. Each pair of individuals decides independently. Indeed, many large real-world graphs, with constant average degree, have a giant component. This is perhaps the most important global property of the $G(n, p)$ model.

3.1.1 Degree Distribution

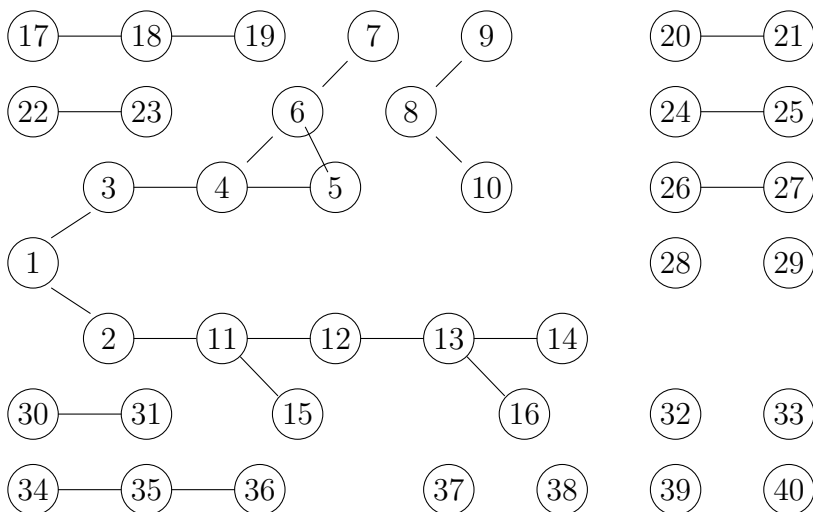
One of the simplest quantities to observe in a real graph is the number of vertices of given degree, called the vertex degree distribution. It is also very simple to study these distributions in $G(n, p)$ since the degree of each vertex is the sum of $n - 1$ independent random variables, which results in a binomial distribution. Since we will be dealing with graphs where the number of vertices n , is large, from here on we often replace $n - 1$ by n to simplify formulas.

Example: In $G(n, \frac{1}{2})$, each vertex is of degree close to $n/2$. In fact, for any $\varepsilon > 0$, the degree of each vertex almost surely is within $1 \pm \varepsilon$ times $n/2$. To see this, note that the probability that a vertex is of degree k is

$$\text{Prob}(k) = \binom{n-1}{k} \left(\frac{1}{2}\right)^k \left(\frac{1}{2}\right)^{n-k} \approx \binom{n}{k} \left(\frac{1}{2}\right)^k \left(\frac{1}{2}\right)^{n-k} = \frac{1}{2^n} \binom{n}{k}.$$



A graph with 40 vertices and 24 edges



A randomly generated $G(n, p)$ graph with 40 vertices and 24 edges

Figure 3.2: Two graphs, each with 40 vertices and 24 edges. The second graph was randomly generated using the $G(n, p)$ model with $p = 1.2/n$. A graph similar to the top graph is almost surely not going to be randomly generated in the $G(n, p)$ model, whereas a graph similar to the lower graph will almost surely occur. Note that the lower graph consists of a giant component along with a number of small components that are trees.

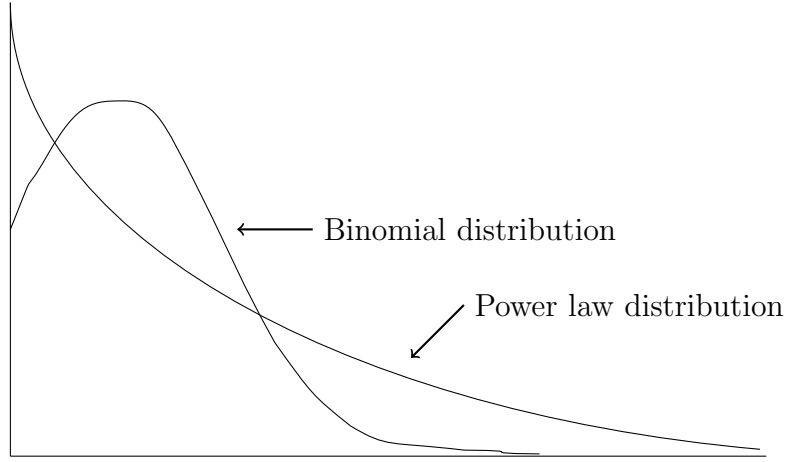


Figure 3.3: Illustration of the binomial and the power law distributions.

This probability distribution has a mean $m = n/2$ and variance $\sigma^2 = n/4$. To see this, observe that the degree k is the sum of n indicator variables that take on value zero or one depending whether an edge is present or not. The expected value of the sum is the sum of the expected values and the variance of the sum is the sum of the variances.

Near the mean, the binomial distribution is well approximated by the normal distribution. See Section 11.4.9 in the appendix.

$$\frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(k-m)^2}{\sigma^2}} = \frac{1}{\sqrt{\pi n/2}} e^{-\frac{(k-n/2)^2}{n/2}}$$

The standard deviation of the normal distribution is $\frac{\sqrt{n}}{2}$ and essentially all of the probability mass is within an additive term $\pm c\sqrt{n}$ of the mean $n/2$ for some constant c and thus is certainly within a multiplicative factor of $1 \pm \varepsilon$ of $n/2$ for sufficiently large n . ■

The degree distribution of $G(n, p)$ for general p is also binomial. Since p is the probability of an edge being present, the expected degree of a vertex is $d \approx pn$. The actual degree distribution is given by

$$\text{Prob}(\text{vertex has degree } k) = \binom{n-1}{k} p^k (1-p)^{n-k-1} \approx \binom{n}{k} p^k (1-p)^{n-k}.$$

The quantity $\binom{n-1}{k}$ is the number of ways of choosing k edges, out of the possible $n-1$ edges, and $p^k (1-p)^{n-k-1}$ is the probability that the k selected edges are present and the remaining $n-k-1$ are not. Since n is large, replacing $n-1$ by n does not cause much error.

The binomial distribution falls off exponentially fast as one moves away from the mean. However, the degree distributions of graphs that appear in many applications do not exhibit such sharp drops. Rather, the degree distributions are much broader. This is often

referred to as having a “heavy tail”. The term tail refers to values of a random variable far away from its mean, usually measured in number of standard deviations. Thus, although the $G(n, p)$ model is important mathematically, more complex models are needed to represent real world graphs.

Consider an airline route graph. The graph has a wide range of degrees, from degree one or two for a small city, to degree 100, or more, for a major hub. The degree distribution is not binomial. Many large graphs that arise in various applications appear to have power law degree distributions. A power law degree distribution is one in which the number of vertices having a given degree decreases as a power of the degree, as in

$$\text{Number}(\text{degree } k \text{ vertices}) = c \frac{n}{k^r},$$

for some small positive real r , often just slightly less than three. Later, we will consider a random graph model giving rise to such degree distributions.

The following theorem claims that the degree distribution of the random graph $G(n, p)$ is tightly concentrated about its expected value. That is, the probability that the degree of a vertex differs from its expected degree, np , by more than $\lambda\sqrt{np}$, drops off exponentially fast with λ .

Theorem 3.1 *Let v be a vertex of the random graph $G(n, p)$. Let α be a real number in $(0, \sqrt{np})$.*

$$\text{Prob}(|np - \deg(v)| \geq \alpha\sqrt{np}) \leq 3e^{-\alpha^2/8}.$$

Proof: The degree $\deg(v)$ is the sum of $n - 1$ independent Bernoulli random variables, $y_1, y_2, \dots, y_{n-1}$, where, y_i is the indicator variable that the i^{th} edge from v is present. So the theorem follows from Theorem 2.12. ■

Theorem 3.1.1 was for one vertex. The corollary below deals with all vertices.

Corollary 3.2 *Suppose ε is a positive constant. If p is $\Omega(\ln n/n\varepsilon^2)$, then, almost surely, every vertex has degree in the range $(1 - \varepsilon)np$ to $(1 + \varepsilon)np$.*

Proof: Apply Theorem with $\alpha = \varepsilon\sqrt{np}$ to get that the probability that an individual vertex has degree outside the range $[(1 - \varepsilon)np, (1 + \varepsilon)np]$ is at most $3e^{-\varepsilon^2 np/8}$. By the union bound, the probability that some vertex has degree outside this range is at most $3ne^{-\varepsilon^2 np/8}$. For this to be $o(1)$, it suffices for p to be $\Omega(\ln n/n\varepsilon^2)$. Hence the Corollary. ■

Note that the assumption p is $\Omega(\ln n/n\varepsilon^2)$ is necessary. If $p = d/n$ for d a constant, then, indeed, some vertices may have degrees outside the range. Without the $\Omega(\ln n/n\varepsilon^2)$ assumption, for $p = \frac{1}{n}$, Corollary 3.1.1 would claim almost surely no vertex had a degree that was greater than a constant independent of n . But shortly we will see that it is highly

likely that for $p = \frac{1}{n}$ there is a vertex of degree $\Omega(\log n / \log \log n)$.

When p is a constant, the expected degree of vertices in $G(n, p)$ increases with n . For example, in $G(n, \frac{1}{2})$, the expected degree of a vertex is $n/2$. In many real applications, we will be concerned with $G(n, p)$ where $p = d/n$, for d a constant, i.e., graphs whose expected degree is a constant d independent of n . Holding $d = np$ constant as n goes to infinity, the binomial distribution

$$\text{Prob}(k) = \binom{n}{k} p^k (1-p)^{n-k}$$

approaches the Poisson distribution

$$\text{Prob}(k) = \frac{(np)^k}{k!} e^{-np} = \frac{d^k}{k!} e^{-d}.$$

To see this, assume $k = o(n)$ and use the approximations $n - k \cong n$, $\binom{n}{k} \cong \frac{n^k}{k!}$, and $(1 - \frac{1}{n})^{n-k} \cong e^{-1}$ to approximate the binomial distribution by

$$\lim_{n \rightarrow \infty} \binom{n}{k} p^k (1-p)^{n-k} = \frac{n^k}{k!} \left(\frac{d}{n}\right)^k \left(1 - \frac{d}{n}\right)^n = \frac{d^k}{k!} e^{-d}.$$

Note that for $p = \frac{d}{n}$, where d is a constant independent of n , the probability of the binomial distribution falls off rapidly for $k > d$, and is essentially zero for all but some finite number of values of k . This justifies the $k = o(n)$ assumption. Thus, the Poisson distribution is a good approximation.

Example: In $G(n, \frac{1}{n})$ many vertices are of degree one, but not all. Some are of degree zero and some are of degree greater than one. In fact, it is highly likely that there is a vertex of degree $\Omega(\log n / \log \log n)$. The probability that a given vertex is of degree k is

$$\text{Prob}(k) = \binom{n}{k} \left(\frac{1}{n}\right)^k \left(1 - \frac{1}{n}\right)^{n-k} \approx \frac{e^{-1}}{k!}.$$

If $k = \log n / \log \log n$,

$$\log k^k = k \log k \cong \frac{\log n}{\log \log n} (\log \log n - \log \log \log n) \cong \log n$$

and thus $k^k \cong n$. Since $k! \leq k^k \cong n$, the probability that a vertex has degree $k = \log n / \log \log n$ is at least $\frac{1}{k!} e^{-1} \geq \frac{1}{en}$. If the degrees of vertices were independent random variables, then this would be enough to argue that there would be a vertex of degree $\log n / \log \log n$ with probability at least $1 - \left(1 - \frac{1}{en}\right)^n = 1 - e^{-\frac{1}{e}} \cong 0.31$. But the degrees are not quite independent since when an edge is added to the graph it affects the degree of two vertices. This is a minor technical point, which one can get around. ■

3.1.2 Existence of Triangles in $G(n, d/n)$

What is the expected number of triangles in $G(n, \frac{d}{n})$, when d is a constant? As the number of vertices increases one might expect the number of triangles to increase, but this is not the case. Although the number of triples of vertices grows as n^3 , the probability of an edge between two specific vertices decreases linearly with n . Thus, the probability of all three edges between the pairs of vertices in a triple of vertices being present goes down as n^{-3} , exactly canceling the rate of growth of triples.

A random graph with n vertices and edge probability d/n , has an expected number of triangles that is independent of n , namely $d^3/6$. There are $\binom{n}{3}$ triples of vertices. Each triple has probability $(\frac{d}{n})^3$ of being a triangle. Let Δ_{ijk} be the indicator variable for the triangle with vertices i, j , and k being present. That is, all three edges (i, j) , (j, k) , and (i, k) being present. Then the number of triangles is $x = \sum_{ijk} \Delta_{ijk}$. Even though the existence of the triangles are not statistically independent events, by linearity of expectation, which does not assume independence of the variables, the expected value of a sum of random variables is the sum of the expected values. Thus, the expected number of triangles is

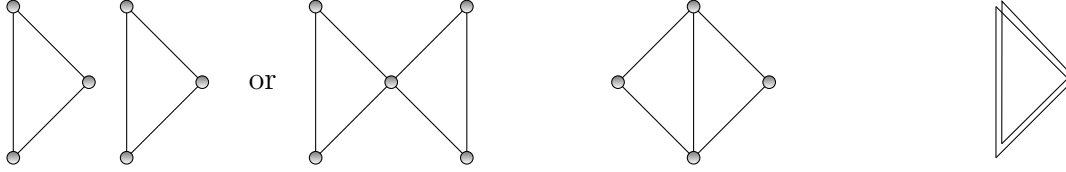
$$E(x) = E\left(\sum_{ijk} \Delta_{ijk}\right) = \sum_{ijk} E(\Delta_{ijk}) = \binom{n}{3} \left(\frac{d}{n}\right)^3 \approx \frac{d^3}{6}.$$

Even though on average there are $\frac{d^3}{6}$ triangles per graph, this does not mean that with high probability a graph has a triangle. Maybe half of the graphs have $\frac{d^3}{3}$ triangles and the other half have none for an average of $\frac{d^3}{6}$ triangles. Then, with probability 1/2, a graph selected at random would have no triangle. If $1/n$ of the graphs had $\frac{d^3}{6}n$ triangles and the remaining graphs had no triangles, then as n goes to infinity, the probability that a graph selected at random would have a triangle would go to zero.

We wish to assert that with some nonzero probability there is at least one triangle in $G(n, p)$ when $p = \frac{d}{n}$. If all the triangles were on a small number of graphs, then the number of triangles in those graphs would far exceed the expected value and hence the variance would be high. A second moment argument rules out this scenario where a small fraction of graphs have a large number of triangles and the remaining graphs have none.

Calculate $E(x^2)$ where x is the number of triangles. Write x as $x = \sum_{ijk} \Delta_{ijk}$, where Δ_{ijk} is the indicator variable of the triangle with vertices i, j , and k being present. Expanding the squared term

$$E(x^2) = E\left(\sum_{i,j,k} \Delta_{ijk}\right)^2 = E\left(\sum_{\substack{i,j,k \\ i',j',k'}} \Delta_{ijk} \Delta_{i'j'k'}\right).$$



The two triangles of Part 1 are either disjoint or share at most one vertex

The two triangles of Part 2 share an edge

The two triangles in Part 3 are the same triangle

Figure 3.4: The triangles in Part 1, Part 2, and Part 3 of the second moment argument for the existence of triangles in $G(n, \frac{d}{n})$.

Split the above sum into three parts. Split the above sum into three parts. In Part 1, let S_1 be the set of i, j, k and i', j', k' which share at most one vertex and hence the two triangles share no edge. In this case, Δ_{ijk} and $\Delta_{i'j'k'}$ are independent and

$$E \left(\sum_{S_1} \Delta_{ijk} \Delta_{i'j'k'} \right) = \sum_{S_1} E(\Delta_{ijk}) E(\Delta_{i'j'k'}) \leq \left(\sum_{\substack{\text{all} \\ ijk}} E(\Delta_{ijk}) \right) \left(\sum_{\substack{\text{all} \\ i'j'k'}} E(\Delta_{i'j'k'}) \right) = E^2(x).$$

In Part 2, i, j, k and i', j', k' share two vertices and hence one edge. See Figure 3.4. Four vertices and five edges are involved overall. There are at most $\binom{n}{4} \in O(n^4)$, 4-vertex subsets and $\binom{4}{2}$ ways to partition the four vertices into two triangles with a common edge. The probability of all five edges in the two triangles being present is p^5 , so this part sums to $O(n^4 p^5) = O(d^5/n)$ and is $o(1)$. In Part 3, i, j, k and i', j', k' are the same sets. The contribution of this part of the summation to $E(x^2)$ is $\binom{n}{3} p^3 = \frac{d^3}{6}$. Thus,

$$E(x^2) \leq E^2(x) + \frac{d^3}{6} + o(1),$$

which implies

$$\text{Var}(x) = E(x^2) - E^2(x) \leq \frac{d^3}{6} + o(1).$$

For x to be less than or equal to zero, it must differ from its expected value by at least its expected value. Thus,

$$\text{Prob}(x = 0) \leq \text{Prob}(|x - E(x)| \geq E(x)).$$

By Chebychev inequality,

$$\text{Prob}(x = 0) \leq \frac{\text{Var}(x)}{E^2(x)} \leq \frac{d^3/6 + o(1)}{d^6/36} \leq \frac{6}{d^3} + o(1). \quad (3.1)$$

Thus, for $d > \sqrt[3]{6} \cong 1.8$, $\text{Prob}(x = 0) < 1$ and $G(n, p)$ has a triangle with nonzero probability. For $d < \sqrt[3]{6}$ and very close to zero, there simply are not enough edges in the graph for there to be a triangle.

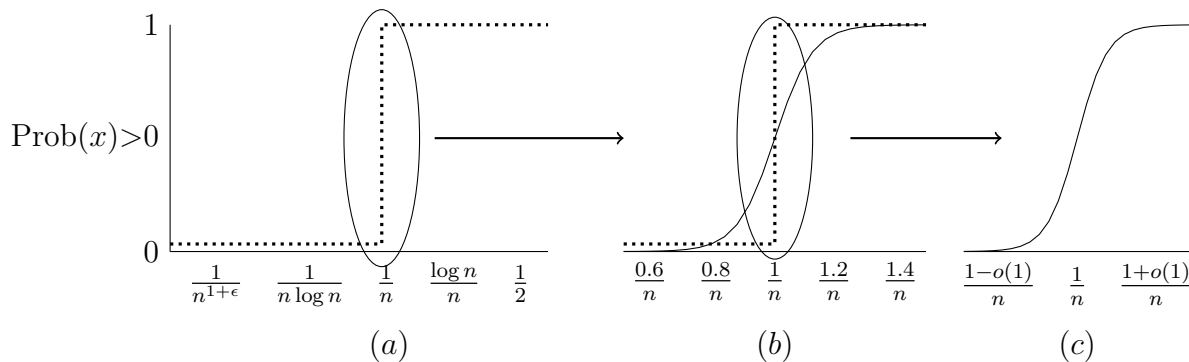


Figure 3.5: Figure 3.5(a) shows a phase transition at $p = \frac{1}{n}$. The dotted line shows an abrupt transition in $\text{Prob}(x) > 0$ from 0 to 1. For any function asymptotically less than $\frac{1}{n}$, $\text{Prob}(x) > 0$ is zero and for any function asymptotically greater than $\frac{1}{n}$, $\text{Prob}(x) > 0$ is one. Figure 3.5(b) expands the scale and shows a less abrupt change in probability unless the phase transition is sharp as illustrated by the dotted line. Figure 3.5(c) is a further expansion and the sharp transition is now more smooth.

3.2 Phase Transitions

Many properties of random graphs undergo structural changes as the edge probability passes some threshold value. This phenomenon is similar to the abrupt phase transitions in physics, as the temperature or pressure increases. Some examples of this are the abrupt appearance of cycles in $G(n, p)$ when p reaches $1/n$ and the disappearance of isolated vertices when p reaches $\frac{\log n}{n}$. The most important of these transitions is the emergence of a giant component, a connected component of size $\Theta(n)$, which happens at $d = 1$. Recall Figure 3.1.

For these and many other properties of random graphs, a threshold exists where an abrupt transition from not having the property to having the property occurs. If there exists a function $p(n)$ such that when $\lim_{n \rightarrow \infty} \frac{p_1(n)}{p(n)} = 0$, $G(n, p_1(n))$ almost surely does not have the property, and when $\lim_{n \rightarrow \infty} \frac{p_2(n)}{p(n)} = \infty$, $G(n, p_2(n))$ almost surely has the property, then we say that a *phase transition* occurs, and $p(n)$ is the *threshold*. Recall that $G(n, p)$ “almost surely does not have the property” means that the probability that it has the property goes to zero in the limit, as n goes to infinity. We shall soon see that every increasing property has a threshold. This is true not only for increasing properties of $G(n, p)$, but for increasing properties of any combinatorial structure. If for $cp(n)$, $c < 1$, the graph almost surely does not have the property and for $cp(n)$, $c > 1$, the graph almost surely has the property, then $p(n)$ is a *sharp threshold*. The existence of a giant component has a sharp threshold at $1/n$. We will prove this later.

In establishing phase transitions, we often use a variable $x(n)$ to denote the number

No items	At least one occurrence of item in 10% of the graphs
$E(x) \geq 0.1$	For 10% of the graphs, $x \geq 1$

Figure 3.6: If the expected fraction of the number of graphs in which an item occurs did not go to zero, then $E(x)$, the expected number of items per graph, could not be zero. Suppose 10% of the graphs had at least one occurrence of the item. Then the expected number of occurrences per graph must be at least 0.1. Thus, $E(x) = 0$ implies the probability that a graph has an occurrence of the item goes to zero. However, the other direction needs more work. If $E(x)$ were not zero, a second moment argument is needed to conclude that the probability that a graph picked at random had an occurrence of the item was nonzero since there could be a large number of occurrences concentrated on a vanishingly small fraction of all graphs. The second moment argument claims that for a nonnegative random variable x with $E(x) > 0$, if $\text{Var}(x)$ is $o(E^2(x))$ or alternatively if $E(x^2) \leq E^2(x)(1 + o(1))$, then almost surely $x > 0$.

of occurrences of an item in a random graph. If the expected value of $x(n)$ goes to zero as n goes to infinity, then a graph picked at random almost surely has no occurrence of the item. This follows from Markov's inequality. Since x is a nonnegative random variable $\text{Prob}(x \geq a) \leq \frac{1}{a}E(x)$, which implies that the probability of $x(n) \geq 1$ is at most $E(x(n))$. That is, if the expected number of occurrences of an item in a graph goes to zero, the probability that there are one or more occurrences of the item in a randomly selected graph goes to zero. This is called the *first moment method*.

The previous section showed that the property of having a triangle has a threshold at $p(n) = 1/n$. If the edge probability $p_1(n)$ is $o(1/n)$, then the expected number of triangles goes to zero and by the first moment method, the graph almost surely has no triangle. However, if the edge probability $p_2(n)$ satisfies $np_2(n) \rightarrow \infty$, then from (3.1), the probability of having no triangle is at most $6/d^3 + o(1) = 6/(np_2(n))^3 + o(1)$, which goes to zero. This latter case uses what we call the second moment method. The first and second moment methods are broadly used. We describe the second moment method in some generality now.

When the expected value of $x(n)$, the number of occurrences of an item, goes to infinity, we cannot conclude that a graph picked at random will likely have a copy since the items may all appear on a small fraction of the graphs. We resort to a technique called

the *second moment method*. It is a simple idea based on Chebyshev's inequality.

Theorem 3.3 (Second Moment method) *Let $x(n)$ be a random variable with $E(x) > 0$. If*

$$\text{Var}(x) = o\left(E^2(x)\right),$$

then x is almost surely greater than zero.

Proof: If $E(x) > 0$, then for x to be less than or equal to zero, it must differ from its expected value by at least its expected value. Thus,

$$\text{Prob}(x \leq 0) \leq \text{Prob}\left(|x - E(x)| \geq E(x)\right).$$

By Chebyshev inequality

$$\text{Prob}\left(|x - E(x)| \geq E(x)\right) \leq \frac{\text{Var}(x)}{E^2(x)} \rightarrow 0.$$

Thus, $\text{Prob}(x \leq 0)$ goes to zero if $\text{Var}(x)$ is $o(E^2(x))$. ■

Corollary 3.4 *Let x be a random variable with $E(x) > 0$. If*

$$E(x^2) \leq E^2(x)(1 + o(1)),$$

then x is almost surely greater than zero.

Proof: If $E(x^2) \leq E^2(x)(1 + o(1))$, then

$$\text{Var}(x) = E(x^2) - E^2(x) \leq E^2(x)o(1) = o(E^2(x)).$$
■

Threshold for graph diameter two

We now present the first example of a sharp phase transition for a property. This means that slightly increasing the edge probability p near the threshold takes us from almost surely not having the property to almost surely having it. The property is that of a random graph having diameter less than or equal to two. The diameter of a graph is the maximum length of the shortest path between a pair of nodes.

The following technique for deriving the threshold for a graph having diameter two is a standard method often used to determine the threshold for many other objects. Let x be a random variable for the number of objects such as triangles, isolated vertices, or Hamilton circuits, for which we wish to determine a threshold. Then we determine the value of p , say p_0 , where the expected value of x goes from zero to infinity. For $p < p_0$

almost surely a graph selected at random will not have a copy of x . For $p > p_0$, a second moment argument is needed to establish that the items are not concentrated on a vanishingly small fraction of the graphs and that a graph picked at random will almost surely have a copy.

Our first task is to figure out what to count to determine the threshold for a graph having diameter two. A graph has diameter two if and only if for each pair of vertices i and j , either there is an edge between them or there is another vertex k to which both i and j have an edge. The set of neighbors of i and the set of neighbors of j are random subsets of expected cardinality np . For these two sets to intersect requires $np \approx \sqrt{n}$ or $p \approx \frac{1}{\sqrt{n}}$. Such statements often go under the general name of “birthday paradox” though it is not a paradox. In what follows, we will prove a threshold of $O(\sqrt{\ln n}/\sqrt{n})$ for a graph to have diameter two. The extra factor of $\sqrt{\ln n}$ ensures that every one of the $\binom{n}{2}$ pairs of i and j has a common neighbor. When $p = c\sqrt{\frac{\ln n}{n}}$, for $c < \sqrt{2}$, the graph almost surely has diameter greater than two and for $c > \sqrt{2}$, the graph almost surely has diameter less than or equal to two.

Theorem 3.5 *The property that $G(n, p)$ has diameter two has a sharp threshold at $p = \sqrt{2}\sqrt{\frac{\ln n}{n}}$.*

Proof: If G has diameter greater than two, then there exists a pair of nonadjacent vertices i and j such that no other vertex of G is adjacent to both i and j . This motivates calling such a pair *bad*.

Introduce a set of indicator random variables I_{ij} , one for each pair of vertices (i, j) with $i < j$, where I_{ij} is 1 if and only if the pair (i, j) is bad. Let

$$x = \sum_{i < j} I_{ij}$$

be the number of bad pairs of vertices. Putting $i < j$ in the sum ensures each pair (i, j) is counted only once. A graph has diameter at most two if and only if it has no bad pair, i.e., $x = 0$. Thus, if $\lim_{n \rightarrow \infty} E(x) = 0$, then for large n , almost surely, a graph has no bad pair and hence has diameter at most two.

The probability that a given vertex is adjacent to both vertices in a pair of vertices (i, j) is p^2 . Hence, the probability that the vertex is not adjacent to both vertices is $1 - p^2$. The probability that no vertex is adjacent to the pair (i, j) is $(1 - p^2)^{n-2}$ and the probability that i and j are not adjacent is $1 - p$. Since there are $\binom{n}{2}$ pairs of vertices, the expected number of bad pairs is

$$E(x) = \binom{n}{2} (1 - p) (1 - p^2)^{n-2}.$$

Setting $p = c\sqrt{\frac{\ln n}{n}}$,

$$\begin{aligned} E(x) &\cong \frac{n^2}{2} \left(1 - c\sqrt{\frac{\ln n}{n}}\right) \left(1 - c^2 \frac{\ln n}{n}\right)^n \\ &\cong \frac{n^2}{2} e^{-c^2 \ln n} \\ &\cong \frac{1}{2} n^{2-c^2}. \end{aligned}$$

For $c > \sqrt{2}$, $\lim_{n \rightarrow \infty} E(x) \rightarrow 0$. Thus, by the first moment method, for $p = c\sqrt{\frac{\ln n}{n}}$ with $c > \sqrt{2}$, $G(n, p)$ almost surely has no bad pair and hence has diameter at most two.

Next, consider the case $c < \sqrt{2}$ where $\lim_{n \rightarrow \infty} E(x) \rightarrow \infty$. We appeal to a second moment argument to claim that almost surely a graph has a bad pair and thus has diameter greater than two.

$$E(x^2) = E\left(\sum_{i < j} I_{ij}\right)^2 = E\left(\sum_{i < j} I_{ij} \sum_{k < l} I_{kl}\right) = E\left(\sum_{\substack{i < j \\ k < l}} I_{ij} I_{kl}\right) = \sum_{\substack{i < j \\ k < l}} E(I_{ij} I_{kl}).$$

The summation can be partitioned into three summations depending on the number of distinct indices among i, j, k , and l . Call this number a .

$$\begin{aligned} E(x^2) &= \sum_{\substack{i < j \\ k < l}} E(I_{ij} I_{kl}) + \sum_{\substack{i < j \\ i < k}} E(I_{ij} I_{ik}) + \sum_{i < j} E(I_{ij}^2). \end{aligned} \tag{3.2}$$

$a = 4 \qquad \qquad a = 3 \qquad \qquad a = 2$

Consider the case $a = 4$ where i, j, k , and l are all distinct. If $I_{ij} I_{kl} = 1$, then both pairs (i, j) and (k, l) are bad and so for each $u \notin \{i, j, k, l\}$, one of the edges (i, u) or (j, u) is absent and, in addition, one of the edges (k, u) or (l, u) is absent. The probability of this for one u not in $\{i, j, k, l\}$ is $(1 - p^2)^2$. As u ranges over all the $n - 4$ vertices not in $\{i, j, k, l\}$, these events are all independent. Thus,

$$E(I_{ij} I_{kl}) \leq (1 - p^2)^{2(n-4)} \leq (1 - c^2 \frac{\ln n}{n})^{2n} (1 + o(1)) \leq n^{-2c^2} (1 + o(1))$$

and the first sum is

$$\sum_{\substack{i < j \\ k < l}} E(I_{ij} I_{kl}) \leq n^{4-2c^2} (1 + o(1)).$$

For the second summation, observe that if $I_{ij} I_{ik} = 1$, then for every vertex u not equal to i, j , or k , either there is no edge between i and u or there is an edge (i, u) and both edges (j, u) and (k, u) are absent. The probability of this event for one u is

$$1 - p + p(1 - p)^2 = 1 - 2p^2 + p^3 \approx 1 - 2p^2.$$

Thus, the probability for all such u is $(1 - 2p^2)^{n-3}$. Substituting $c\sqrt{\frac{\ln n}{n}}$ for p yields

$$\left(1 - \frac{2c^2 \ln n}{n}\right)^{n-3} \cong e^{-2c^2 \ln n} = n^{-2c^2},$$

which is an upper bound on $E(I_{ij}I_{kl})$ for one i, j, k , and l with $a = 3$. Summing over all distinct triples yields n^{3-2c^2} for the second summation in (3.2).

For the third summation, since the value of I_{ij} is zero or one, $E(I_{ij}^2) = E(I_{ij})$. Thus,

$$\sum_{ij} E(I_{ij}^2) = E(x).$$

Hence, $E(x^2) \leq n^{4-2c^2} + n^{3-2c^2} + n^{2-c^2}$ and $E(x) \cong n^{2-c^2}$, from which it follows that for $c < \sqrt{2}$, $E(x^2) \leq E^2(x)(1 + o(1))$. By a second moment argument, Corollary 3.4, a graph almost surely has at least one bad pair of vertices and thus has diameter greater than two. Therefore, the property that the diameter of $G(n, p)$ is less than or equal to two has a sharp threshold at $p = \sqrt{2}\sqrt{\frac{\ln n}{n}}$ ■

Disappearance of Isolated Vertices

The disappearance of isolated vertices in $G(n, p)$ has a sharp threshold at $\frac{\ln n}{n}$. At this point the giant component has absorbed all the small components and with the disappearance of isolated vertices, the graph becomes connected.

Theorem 3.6 *The disappearance of isolated vertices in $G(n, p)$ has a sharp threshold of $\frac{\ln n}{n}$.*

Proof: Let x be the number of isolated vertices in $G(n, p)$. Then,

$$E(x) = n(1 - p)^{n-1}.$$

Since we believe the threshold to be $\frac{\ln n}{n}$, consider $p = c\frac{\ln n}{n}$. Then,

$$\lim_{n \rightarrow \infty} E(x) = \lim_{n \rightarrow \infty} n \left(1 - \frac{c \ln n}{n}\right)^n = \lim_{n \rightarrow \infty} n e^{-c \ln n} = \lim_{n \rightarrow \infty} n^{1-c}.$$

If $c > 1$, the expected number of isolated vertices, goes to zero. If $c < 1$, the expected number of isolated vertices goes to infinity. If the expected number of isolated vertices goes to zero, it follows that almost all graphs have no isolated vertices. On the other hand, if the expected number of isolated vertices goes to infinity, a second moment argument is needed to show that almost all graphs have an isolated vertex and that the isolated vertices are not concentrated on some vanishingly small set of graphs with almost all graphs not having isolated vertices.

Assume $c < 1$. Write $x = I_1 + I_2 + \cdots + I_n$ where I_i is the indicator variable indicating whether vertex i is an isolated vertex. Then $E(x^2) = \sum_{i=1}^n E(I_i^2) + 2 \sum_{i \leq j} E(I_i I_j)$. Since I_i equals 0 or 1, $I_i^2 = I_i$ and the first sum has value $E(x)$. Since all elements in the second sum are equal

$$\begin{aligned} E(x^2) &= E(x) + n(n-1)E(I_1 I_2) \\ &= E(x) + n(n-1)(1-p)^{2(n-1)-1}. \end{aligned}$$

The minus one in the exponent $2(n-1) - 1$ avoids counting the edge from vertex 1 to vertex 2 twice. Now,

$$\begin{aligned} \frac{E(x^2)}{E^2(x)} &= \frac{n(1-p)^{n-1} + n(n-1)(1-p)^{2(n-1)-1}}{n^2(1-p)^{2(n-1)}} \\ &= \frac{1}{n(1-p)^{n-1}} + \left(1 - \frac{1}{n}\right) \frac{1}{1-p}. \end{aligned}$$

For $p = c \frac{\ln n}{n}$ with $c < 1$, $\lim_{n \rightarrow \infty} E(x) = \infty$ and

$$\lim_{n \rightarrow \infty} \frac{E(x^2)}{E^2(x)} = \lim_{n \rightarrow \infty} \left[\frac{1}{n^{1-c}} + \left(1 - \frac{1}{n}\right) \frac{1}{1 - c \frac{\ln n}{n}} \right] = 1 + o(1).$$

By the second moment argument, Corollary 3.4, the probability that $x = 0$ goes to zero implying that almost all graphs have an isolated vertex. Thus, $\frac{\ln n}{n}$ is a sharp threshold for the disappearance of isolated vertices. For $p = c \frac{\ln n}{n}$, when $c > 1$ there almost surely are no isolated vertices, and when $c < 1$ there almost surely are isolated vertices. ■

Hamilton circuits

So far in establishing phase transitions in the $G(n, p)$ model for an item such as the disappearance of isolated vertices, we introduced a random variable x that was the number of occurrences of the item. We then determined the probability p for which the expected value of x went from zero to infinity. For values of p for which $E(x) = 0$, we argued that with probability one, a graph generated at random had no occurrences of x . For values of x for which $E(x) \rightarrow \infty$, we used the second moment argument to conclude that with probability one a graph generated at random had occurrences of x . That is, the occurrences that forced $E(x)$ to infinity were not all concentrated on a vanishingly small fraction of the graphs. One might raise the question for the $G(n, p)$ graph model, do there exist items that are so concentrated on a small fraction of the graphs that the value of p where $E(x)$ goes from zero to infinity is not the threshold? An example where this happens is Hamilton circuits.

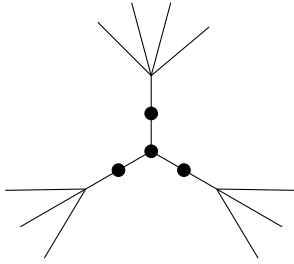


Figure 3.7: A degree three vertex with three adjacent degree two vertices. Graph cannot have a Hamilton circuit.

Let x be the number of Hamilton circuits in $G(n, p)$ and let $p = \frac{d}{n}$ for some constant d . There are $\frac{1}{2}(n-1)!$ potential Hamilton circuits in a graph and each has probability $(\frac{d}{n})^n$ of actually being a Hamilton circuit. Thus,

$$\begin{aligned} E(x) &= \frac{1}{2}(n-1)! \left(\frac{d}{n}\right)^n \\ &\simeq \left(\frac{n}{e}\right)^n \left(\frac{d}{n}\right)^n \\ &= \begin{cases} 0 & d < e \\ \infty & d > e \end{cases} \end{aligned}$$

This suggests that the threshold for Hamilton circuits occurs when d equals Euler's constant e . This is not possible since the graph still has isolated vertices and is not even connected for $p = \frac{e}{n}$. Thus, the second moment argument is indeed necessary.

The actual threshold for Hamilton circuits is $d = \omega(\log n + \log \log n)$. For any $p(n)$ asymptotically greater than $\log n + \log \log n$, $G(n, p)$ will have a Hamilton circuit with probability one. This is the same threshold as for the disappearance of degree one vertices. Clearly a graph with a degree one vertex cannot have a Hamilton circuit. But it may seem surprising that Hamilton circuits appear as soon as degree one vertices disappear. You may ask why at the moment degree one vertices disappear there cannot be a subgraph consisting of a degree three vertex adjacent to three degree two vertices as shown in Figure 3.7. The reason is that the frequency of degree two and three vertices in the graph is very small and the probability that four such vertices would occur together in such a subgraph is too small for it to happen.

3.3 The Giant Component

Consider $G(n, p)$ as p grows. Starting with $p = 0$, the graph has n vertices and no edges. As p increases and edges are added, a forest of trees emerges. When p is $o(1/n)$ the graph is almost surely a forest of trees, i.e., there are no cycles. When p is d/n , d a constant, cycles appear. For $d < 1$, no connected component has asymptotically more

than $\log n$ vertices. The number of components containing a single cycle is a constant independent of n . Thus, the graph consists of a forest of trees plus a few components that have a single cycle with no $\Omega(\log n)$ size components.

At p equal $1/n$, a phase transition occurs in which a giant component emerges. The transition consists of a double jump. At $p = 1/n$, components of $n^{2/3}$ vertices emerge, which are almost surely trees. Then at $p = d/n$, $d > 1$, a true giant component emerges that has a number of vertices proportional to n . This is a seminal result in random graph theory and the main subject of this section. Giant components also arise in many real world graphs; the reader may want to look at large real-world graphs, like portions of the web and find the size of the largest connected component.

When one looks at the connected components of large graphs that appear in various contexts, one observes that often there is one very large component. One example is a graph formed from a data base of protean interactions¹ where vertices correspond to proteins and edges correspond to pairs of proteins that interact. By an interaction, one means two amino acid chains that bind to each other for a function. The graph has 2735 vertices and 3602 edges. At the time we looked at the data base, the associated graph had the number of components of various sizes shown in Table 3.1. There are a number of small components, but only one component of size greater than 16, and that is a giant component of size 1851. As more proteins are added to the data base the giant component will grow even larger and eventually swallow up all the smaller components.

Size of component	1	2	3	4	5	6	7	8	9	10	11	12	...	15	16	...	1851
Number of components	48	179	50	25	14	6	4	6	1	1	1	0	0	0	1	0	1

Table 1: Table 3.1 Size of components in the graph implicit in the database of interacting proteins.

The existence of a giant component is not unique to the graph produced from the protein data set. Take any data set that one can convert to a graph and it is likely that the graph will have a giant component, provided that the ratio of edges to vertices is a small number greater than one half. Table 3.2 gives two other examples. This phenomenon, of the existence of a giant component in many real world graphs deserves study.

Returning to $G(n, p)$, as p increases beyond d/n , all nonisolated vertices are absorbed into the giant component, and at $p = \frac{1}{2} \frac{\ln n}{n}$, the graph consists only of isolated vertices plus a giant component. At $p = \frac{\ln n}{n}$, the graph becomes completely connected. By $p = 1/2$, the graph is not only connected, but is sufficiently dense that it has a clique of

¹Science 1999 July 30 Vol. 285 No. 5428 pp751-753.

<ftp://ftp.cs.rochester.edu/pub/u/joel/papers.lst>

Vertices are papers and edges mean that two papers shared an author.

1	2	3	4	5	6	7	8	14	27488
2712	549	129	51	16	12	8	3	1	1

<http://www.gutenberg.org/etext/3202>

Vertices represent words and edges connect words that are synonyms of one another.

1	2	3	4	5	14	16	18	48	117	125	128	30242
7	1	1	1	0	1	1	1	1	1	1	1	1

Table 2: Table 3.2 Size of components in two graphs constructed from data sets.

size $(2 - \varepsilon) \log n$ for any $\varepsilon > 0$. We prove many of these facts in this chapter.

To compute the size of a connected component of $G(n, p)$, do a breadth first search of a component starting from an arbitrary vertex and generate an edge only when the search process needs to know if the edge exists. Start at an arbitrary vertex and mark it discovered and unexplored. At a general step, select a discovered, but unexplored vertex v , and explore it as follows. For each undiscovered vertex u , independently decide with probability $p = d/n$ whether the edge (v, u) is in and if it is, mark u discovered and unexplored. After this, mark v explored. Discovered but unexplored vertices are called the frontier. The algorithm has found the entire connected component when the frontier becomes empty.

For each vertex u , other than the start vertex, the probability that u is undiscovered after the first i steps is precisely $(1 - \frac{d}{n})^i$. A step is the full exploration of one vertex. Let z_i be the number of vertices discovered in the first i steps of the search. The distribution of z_i is Binomial $\left(n - 1, 1 - (1 - \frac{d}{n})^i\right)$.

Consider the case $d > 1$. For small values of i , the probability that a vertex is undiscovered after i steps is

$$\left(1 - \frac{d}{n}\right)^i \approx 1 - \frac{id}{n}.$$

The probability that a vertex is discovered after i steps is $\frac{id}{n}$. The expected number of discovered vertices grows as id and the expected size of the frontier grows as $(d - 1)i$. As the fraction of discovered vertices increases, the expected rate of growth of newly discovered vertices decreases since many of the vertices adjacent to the vertex currently being searched have already been discovered. Once $\frac{d-1}{d}n$ vertices have been discovered, the growth of newly discovered vertices slows to one at each step. Eventually for $d > 1$, the growth of discovering new vertices drops below one per step and the frontier starts

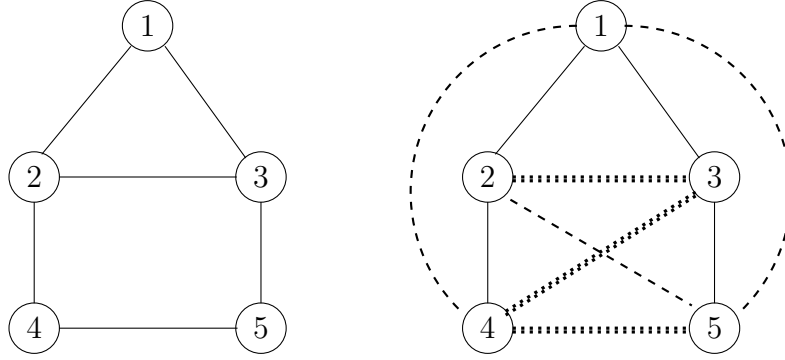


Figure 3.8: A graph (left) and the breadth first search of the graph (right). At vertex 1 the algorithm queried all edges. The solid edges are real edges, the dashed edges are edges that were queried but do not exist. At vertex 2 the algorithm queried all possible edges to vertices not yet discovered. The algorithm does not query whether the edge (2,3) exists since vertex 3 has already been discovered when the algorithm is at vertex 2. Potential edges not queried are illustrated with dotted edges.

to shrink. For $d < 1$, $(d - 1)i$, the expected size of the frontier is negative. The expected rate of growth is less than one, even at the start.

It is easy to make this argument rigorous to prove that for the $d < 1$ case, almost surely, there is no connected component of size $\Omega(\ln n)$. We do this before tackling the more difficult $d > 1$ case.

Theorem 3.7 *Let $p = d/n$ with $d < 1$. The probability that $G(n, p)$ has a component of size more than $c \frac{\ln n}{(1-d)^2}$ is at most $1/n$ for a suitable constant c depending on d but not on n .*

Proof: There is a connected component of size at least k containing a particular vertex v only if the breadth first search started at v has a nonempty frontier at all times up to k . Let z_k be the number of discovered vertices after k steps. The probability that v is in a connected component of size greater than or equal to k is less than or equal to $\text{Prob}(z_k > k)$. Now the distribution of z_k is $\text{Binomial}(n - 1, 1 - (1 - d/n)^k)$. Since $(1 - d/n)^k \geq 1 - dk/n$, the mean of $\text{Binomial}(n - 1, 1 - (1 - d/n)^k)$ is less than the mean of $\text{Binomial}(n, \frac{dk}{n})$. Since $\text{Binomial}(n, \frac{dk}{n})$ has mean dk , the mean of z_k is at most dk where $d < 1$. By a Chernoff bound, the probability that z_k is greater than k is at most $e^{-c_0 k}$ for some constant $c_0 > 0$. If $k \geq c \ln n$ for a suitably large c , then this probability is at most $1/n^2$. This bound is for a single vertex v . Multiplying by n for the union bound completes the proof. ■

Now assume $d > 1$. As we saw, the expected size of the frontier grows as $(d - 1)i$ for small i . The actual size of the frontier is a random variable. What is the probability that the actual size of the frontier will differ from the expected size of the frontier by

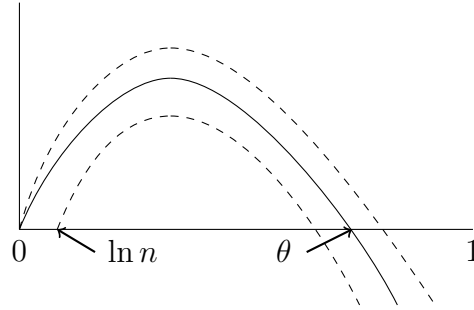


Figure 3.9: The solid curve is the expected size of the frontier. The two dashed curves indicate the range of possible values for the actual size of the frontier.

a sufficient amount so that the actual size of the frontier is zero? To answer this, we need to understand the distribution of the number of discovered vertices after i steps. For small i , the probability that a vertex has been discovered is $1 - (1 - d/n)^i \approx id/n$ and the binomial distribution for the number of discovered vertices, $\text{binomial}(n, \frac{id}{n})$, is well approximated by the Poisson distribution with the same mean id . The probability that a total of k vertices have been discovered in i steps is approximately $e^{-di} \frac{(di)^k}{k!}$. For a connected component to have exactly i vertices, the frontier must drop to zero for the first time at step i . A necessary condition is that exactly i vertices must have been discovered in the first i steps. The probability of this approximately equals

$$e^{-di} \frac{(di)^i}{i!} = e^{-di} \frac{d^i i^i}{i!} e^i = e^{-(d-1)i} d^i = e^{-(d-1-\ln d)i}.$$

For $d > 1$, $\ln d \leq d-1$ and hence $d-1-\ln d > 0$. This probability drops off exponentially with i . For $i > c \ln n$ and sufficiently large c , the probability that the breadth first search starting from a particular vertex terminates with a component of size i is $o(1/n)$ as long as the Poisson approximation is valid. In the range of this approximation, the probability that a breadth first search started from any vertex terminates with $i > c \ln n$ vertices is $o(1)$. Intuitively, if the component has not stopped growing within $\Omega(\ln n)$ steps, it is likely to continue to grow until it becomes much larger and the expected value of the size of the frontier again becomes small. While the expected value of the frontier is large, the probability that the actual size will differ from the expected size sufficiently for the actual size of the frontier to be zero is vanishingly small.

In Theorem 3.9, we prove that there is one giant component of size $\Omega(n)$ along with a number of components of size $O(\ln n)$. We first prove a technical lemma stating that the probability of a vertex being in a small component is strictly less than one and hence there is a giant component. We refer to a connected component of size $O(\log n)$ as a small component.

Lemma 3.8 *Assume $d > 1$. The probability that $cc(v)$, the connected component containing vertex v , is small (i.e., of size $O(\log n)$) is a constant strictly less than 1.*

Proof: Let p be the probability that $cc(v)$ is small, i.e., the probability that the breadth first search started at v terminates before $c_1 \log n$ vertices are discovered. Slightly modify the breadth first search as follows: If in exploring a vertex u at some point, there are m undiscovered vertices, choose the number k of vertices which will be adjacent to u from $\text{Binomial}(m, \frac{d}{n})$ distribution. Having picked k , pick one of the $\binom{m}{k}$ subsets of m undiscovered vertices to be the set of vertices adjacent to u , and make the other $m - k$ vertices not adjacent to u . This process has the same distribution as picking each edge from u independently at random to be present with probability d/n . As the search proceeds, m decreases. If $cc(v)$ is small, m is always greater than $s = n - c_1 \log n$. Modify the process once more picking k from $\text{Binomial}(s, \frac{d}{n})$ instead of from $\text{Binomial}(m, \frac{d}{n})$. Let p' be the probability that $cc(v)$ is small for the modified process. Clearly, $p' \geq p$, so it suffices to prove that p' is a constant strictly less than one. The mean of the binomial now is $d_1 = sd/n$ which is strictly greater than one. It is clear that the probability that the modified process ends before $c_1 \log n$ vertices are discovered is at least the probability for the original process, since picking from $n - c_1 \log n$ vertices has decreased the number of newly discovered vertices each time. Modifying the process so that the newly discovered vertices are picked from a fixed size set, converts the problem to what is called a branching process..

A branching process is a method for creating a possibly infinite random tree. There is a nonnegative integer-valued random variable y that is the number of children of the node being explored. First, the root v of the tree chooses a value of y according to the distribution of y and spawns that number of children. Each of the children independently chooses a value according to the same distribution of y and spawns that many children. The process terminates when all of the vertices have spawned children. The process may go on forever. If it does terminate with a finite tree, we say that the process has become “extinct”. Let $\text{Binomial}(s, \frac{d}{n})$ be the distribution of y . Let q be the probability of extinction. Then, $q \geq p'$, since, the breadth first search terminating with at most $c_1 \log n$ vertices is one way of becoming extinct. Let $p_i = \binom{s}{i} (d/n)^i (1 - (d/n))^{s-i}$ be the probability that y spawns i children. We have $\sum_{i=0}^s p_i = 1$ and $\sum_{i=1}^s i p_i = E(y) = ds/n > 1$.

The depth of a tree is at most the number of nodes in the tree. Let a_t be the probability that the branching process terminates at depth at most t . If the root v has no children, then the process terminates with depth one where the root is counted as a depth one node which is at most t . If v has i children, the process from v terminates at depth at most t if and only if the i sub processes, one rooted at each child of v terminate at depth $t - 1$ or less. The i processes are independent, so the probability that they all terminate at depth at most $t - 1$ is exactly a_{t-1}^i . With this we get:

$$a_t = p_0 + \sum_{i=1}^s p_i a_{t-1}^i = \sum_{i=0}^s p_i a_{t-1}^i.$$

For a small number i of steps, the probability distribution of the size of the set of discovered vertices at time i is $p(k) = e^{-di} \frac{(di)^k}{k!}$ and has expected value di . Thus, the expected size of the frontier is $(d-1)i$. For the frontier to be empty would require that the size of the set of discovered vertices be smaller than its expected value by $(d-1)i$. That is, the size of the set of discovered vertices would need to be $di - (d-1)i = i$. The probability of this is

$$e^{-di} \frac{(di)^i}{i!} = e^{-di} \frac{d^i i^i}{i!} e^i = e^{-(d-1)i} d^i = e^{-(d-1-\ln d)i}$$

which drops off exponentially fast with i provided $d > 1$. Since $d - 1 - \ln d$ is some constant $c > 0$, the probability is e^{-ci} which for $i = \ln n$ is $e^{-c \ln n} = \frac{1}{n^c}$. Thus, with high probability, the largest small component in the graph is of size at most $\ln n$.

Illustration 3.1

We have $a_1 = p_0 < 1$. There is a constant $\alpha \in [p_0, 1)$ such that whenever $a_{t-1} \leq \alpha$, the above recursion implies that $a_t \leq \alpha$. This would finish the proof since then $a_1 \leq \alpha$ implies $a_2 \leq \alpha$ which implies $a_3 \leq \alpha$ etc. and so $q = \lim_{t \rightarrow \infty} a_t \leq \alpha$.

To prove the claim, consider the polynomial

$$h(x) = x - \sum_{i=0}^s p_i x^i.$$

We see that $h(1) = 0$ and $h'(1) = 1 - \sum_{i=1}^s i p_i \approx 1 - \frac{sd}{n}$, which is at most a strictly negative constant. By continuity of $h(\cdot)$, there exists some $x_0 < 1$ such that $h(x) \geq 0$ for $x \in [x_0, 1]$. Take $\alpha = \max(x_0, p_0)$. Now since $\sum_{i=0}^s p_i x^i$ has all nonnegative coefficients, it is an increasing function of x and so if a_{t-1} is at least α , then, $\sum_{i=0}^s p_i a_{t-1}^i$ is at least $\sum_{i=0}^s p_i \alpha^i \geq \alpha$. Now, if $a_{t-1} \leq \alpha$,

$$a_t = \sum_{i=0}^s p_i a_{t-1}^i \geq \sum_{i=1}^s p_i \alpha^i = \alpha - h(\alpha) \leq \alpha,$$

proving the claim. ■

We now prove in Theorem 3.9 that in $G(n, \frac{d}{n})$, $d > 1$ there is one giant component containing a fraction of the n vertices and that the remaining vertices are in components of size less than some constant c_1 times $\log n$. There are no components greater than $c_1 \log n$ other than the giant component.

Theorem 3.9 *Let $p = d/n$ with $d > 1$.*

1. *There are constants c_1 and c_2 such that the probability that there is a connected component of size between $c_1 \log n$ and $c_2 n$ is at most $1/n$.*

2. The number of vertices in components of size $O(\ln n)$ is almost surely at most cn for some $c < 1$. Thus, with probability $1 - o(1)$, there is a connected component of size $\Omega(n)$.
3. The probability that there are two or more connected components, each of size more than $n^{2/3}$, is at most $1/n$.

Proof: In the breadth first search of a component, the probability that a vertex has not been discovered in i steps is $(1 - \frac{d}{n})^i$. It is easy to see that the approximation $(1 - d/n)^i \approx 1 - id/n$ is valid as long as $i \leq c_2 n$ for a suitable constant c_2 since the error term in the approximation is $O(i^2 d^2 / n^2)$, which for $i \leq c_2 n$ is at most a small constant times id/n . This establishes (1).

Next consider (2). For a vertex v , let $\text{cc}(v)$ denote the set of vertices in the connected component containing v . By (1), almost surely, $\text{cc}(v)$ is a small set of size at most $c_1 \log n$ or a large set of size at least $c_2 n$ for every vertex v . The central part of the proof of (2) that the probability of a vertex being in a small component is strictly less than one was established in Lemma 3.8. Let x be the number of vertices in a small connected component. Lemma 3.8 implies that the expectation of the random variable x equals the number of vertices in small connected components is at most some $c_3 n$, for a constant c_3 strictly less than one. But we need to show that for any graph almost surely the actual number x of such vertices is at most some constant strictly less than one times n . For this, we use the second moment method. In this case, the proof that the variance of x is $o(E^2(x))$ is easy. Let x_i be the indicator random variable of the event that $\text{cc}(i)$ is small. Let S and T run over all small sets. Noting that for $i \neq j$, $\text{cc}(i)$ and $\text{cc}(j)$ either are the same or are disjoint,

$$\begin{aligned}
E(x^2) &= E\left(\left(\sum_{i=1}^n x_i\right)^2\right) = \sum_{i,j} E(x_i x_j) = \sum_i E(x_i^2) + \sum_{i \neq j} E(x_i x_j) \\
&= E(x) + \sum_{i \neq j} \sum_S \text{Prob}(\text{cc}(i) = \text{cc}(j) = S) + \sum_{i \neq j} \sum_{\substack{S,T \\ \text{disjoint}}} \text{Prob}(\text{cc}(i) = S; \text{cc}(j) = T) \\
&= E(x) + \sum_{i \neq j} \sum_S \text{Prob}(\text{cc}(i) = \text{cc}(j) = S) \\
&\quad + \sum_{i \neq j} \sum_{\substack{S,T \\ \text{disjoint}}} \text{Prob}(\text{cc}(i) = S) \text{Prob}(\text{cc}(j) = T) (1-p)^{-|S||T|} \\
&\leq O(n) + (1-p)^{-|S||T|} \left(\sum_S \text{Prob}(\text{cc}(i) = S)\right) \left(\sum_T \text{Prob}(\text{cc}(j) = T)\right) \\
&\leq O(n) + (1+o(1)) E(x)E(x).
\end{aligned}$$

In the next to last line, if S containing i and T containing j are disjoint sets, then the two events, S is a connected component and T is a connected component, depend on disjoint

sets of edges except for the $|S||T|$ edges between S vertices and T vertices. Let c_4 be a constant in the interval $(c_3, 1)$. Then, by Chebyshev inequality,

$$\text{Prob}(x > c_4 n) \leq \frac{\text{Var}(x)}{(c_4 - c_3)^2 n^2} \leq \frac{O(n) + o(1)c_3^2 n^2}{(c_4 - c_3)^2 n^2} = o(1).$$

For the proof of (3) suppose a pair of vertices u and v belong to two different connected components, each of size at least $n^{2/3}$. With high probability, they should have merged into one component producing a contradiction. First, run the breadth first search process starting at v for $\frac{1}{2}n^{2/3}$ steps. Since v is in a connected component of size $n^{2/3}$, there are $\Omega(n^{2/3})$ frontier vertices. The expected size of the frontier continues to grow until some constant times n and the actual size of the frontier does not differ significantly from the expected size. The size of the component also grows linearly with n . Thus, the frontier is of size $n^{\frac{2}{3}}$. See Exercise 3.26. By the assumption, u does not belong to this connected component. Now, temporarily stop the breadth first search tree of v and begin a breadth first search tree starting at u , again for $\frac{1}{2}n^{2/3}$ steps. It is important to understand that this change of order of building $G(n, p)$ does not change the resulting graph. We can choose edges in any order since the order does not affect independence or conditioning. The breadth first search tree from u also will have $\Omega(n^{2/3})$ frontier vertices with high probability. Now grow the u tree further. The probability that none of the edges between the two frontier sets is encountered is $(1 - p)^{\Omega(n^{4/3})} \leq e^{-\Omega(dn^{1/3})}$, which converges to zero. So almost surely, one of the edges is encountered and u and v end up in the same connected component. This argument shows for a particular pair of vertices u and v , the probability that they belong to different large connected components is very small. Now use the union bound to conclude that this does not happen for any of the $\binom{n}{2}$ pairs of vertices. The details are left to the reader. ■

3.4 Branching Processes

A *branching process* is a method for creating a random tree. Starting with the root node, each node has a probability distribution for the number of its children. The root of the tree denotes a parent and its descendants are the children with their descendants being the grandchildren. The children of the root are the first generation, their children the second generation, and so on. Branching processes have obvious applications in population studies, but also in exploring a connected component in a random graph.

We analyze a simple case of a branching process where the distribution of the number of children at each node in the tree is the same. The basic question asked is what is the probability that the tree is finite, i.e., the probability that the branching process dies out? This is called the *extinction probability*.

Our analysis of the branching process will give the probability of extinction, as well as the expected size of the components conditioned on extinction. Not surprisingly, the

expected size of components conditioned on extinction is $O(1)$. This says that in $G(n, \frac{d}{n})$, with $d > 1$, there is one giant component of size $\Omega(n)$, the rest of the components are $O(\ln n)$ in size and the expected size of the small components is $O(1)$.

An important tool in our analysis of branching processes is the generating function. The generating function for a nonnegative integer valued random variable y is $f(x) = \sum_{i=0}^{\infty} p_i x^i$ where p_i is the probability that y equals i . The reader not familiar with generating functions should consult Section 11.7 of the appendix.

Let the random variable z_j be the number of children in the j^{th} generation and let $f_j(x)$ be the generating function for z_j . Then $f_1(x) = f(x)$ is the generating function for the first generation where $f(x)$ is the generating function for the number of children at a node in the tree. The generating function for the 2^{nd} generation is $f_2(x) = f(f(x))$. In general, the generating function for the $j+1^{st}$ generation is given by $f_{j+1}(x) = f_j(f(x))$. To see this, observe two things.

First, the generating function for the sum of two identically distributed integer valued random variables x_1 and x_2 is the square of their generating function

$$f^2(x) = p_0^2 + (p_0 p_1 + p_1 p_0)x + (p_0 p_2 + p_1 p_1 + p_2 p_0)x^2 + \cdots$$

For $x_1 + x_2$ to have value zero, both x_1 and x_2 must have value zero, for $x_1 + x_2$ to have value one, exactly one of x_1 or x_2 must have value zero and the other have value one, and so on. In general, the generating function for the sum of i independent random variables, each with generating function $f(x)$, is $f^i(x)$.

The second observation is that the coefficient of x^i in $f_j(x)$ is the probability of there being i children in the j^{th} generation. If there are i children in the j^{th} generation, the number of children in the $j+1^{st}$ generation is the sum of i independent random variables each with generating function $f(x)$. Thus, the generating function for the $j+1^{st}$ generation, given i children in the j^{th} generation, is $f^i(x)$. The generating function for the $j+1^{st}$ generation is given by

$$f_{j+1}(x) = \sum_{i=0}^{\infty} \text{Prob}(z_j = i) f^i(x).$$

If $f_j(x) = \sum_{i=0}^{\infty} a_i x^i$, then f_{j+1} is obtained by substituting $f(x)$ for x in $f_j(x)$.

Since $f(x)$ and its iterates, $f_2, f_3, \dots$, are all polynomials in x with nonnegative coefficients, $f(x)$ and its iterates are all monotonically increasing and convex on the unit interval. Since the probabilities of the number of children of a node sum to one, if $p_0 < 1$, some coefficient of x to a power other than zero in $f(x)$ is nonzero and $f(x)$ is strictly

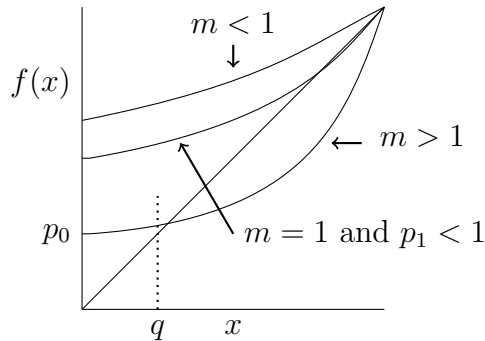


Figure 3.10: Illustration of the root of equation $f(x) = x$ in the interval $[0,1]$.

increasing.

We now focus on roots of the equation $f(x) = x$ in the interval $[0,1]$. The value $x = 1$ is always a root of the equation $f(x) = x$ since $f(1) = \sum_{i=0}^{\infty} p_i = 1$. When is there a smaller nonnegative root? The derivative of $f(x)$ at $x = 1$ is $f'(1) = p_1 + 2p_2 + 3p_3 + \dots$. Let $m = f'(1)$. Thus, m is the expected number of children of a node. If $m > 1$, one might expect the tree to grow forever, since each node at time j is expected to have more than one child. But this does not imply that the probability of extinction is zero. In fact, if $p_0 > 0$, then with positive probability, the root will have no children and the process will become extinct right away. Recall that for $G(n, \frac{d}{n})$, the expected number of children is d , so the parameter m plays the role of d .

If $m < 1$, then the slope of $f(x)$ at $x = 1$ is less than one. This fact along with convexity of $f(x)$ implies that $f(x) > x$ for x in $[0,1]$.

If $m = 1$ and $p_1 < 1$, then once again convexity implies that $f(x) > x$ for $x \in [0,1]$ and there is no root of $f(x)$ in the interval $[0,1]$. If $m = 1$ and $p_1 = 1$, then $f(x)$ is the straight line $f(x) = x$.

If $m > 1$, then the slope of $f(x)$ is greater than the slope of x at $x = 1$. This fact, along with convexity of $f(x)$, implies $f(x) = x$ has a unique root in $[0,1]$. When $p_0 = 0$, the root is at $x = 0$.

Let q be the smallest nonnegative root of the equation $f(x) = x$. For $m < 1$ and for $m=1$ and $p_0 < 1$, q equals one and for $m > 1$, q is strictly less than one. We shall see that the value of q is the *extinction probability* of the branching process and that $1 - q$ is the *immortality probability*. That is, q is the probability that for some j , the number of children in the j^{th} generation is zero. To see this, note that for $m > 1$, $\lim_{j \rightarrow \infty} f_j(x) = q$ for $0 < x < 1$. Figure 3.11 illustrates the proof which is given in Lemma 3.10. Similarly note

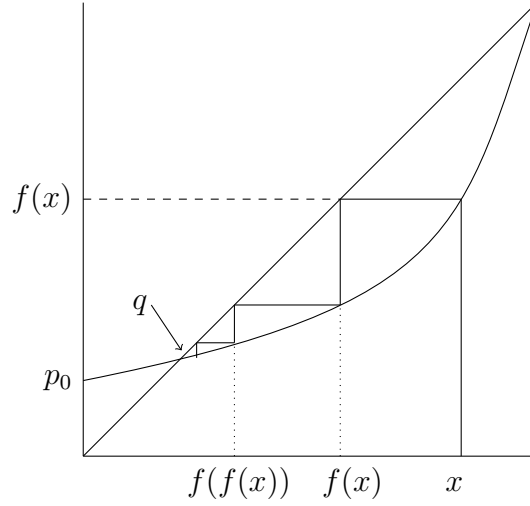


Figure 3.11: Illustration of convergence of the sequence of iterations $f_1(x), f_2(x), \dots$ to q .

that when $m < 1$ or $m = 1$ with $p_0 < 1$, $f_j(x)$ approaches one as j approaches infinity.

Lemma 3.10 Assume $m > 1$. Let q be the unique root of $f(x)=x$ in $[0,1)$. In the limit as j goes to infinity, $f_j(x) = q$ for x in $[0,1)$.

Proof: If $0 \leq x \leq q$, then $x < f(x) \leq f(q)$ and iterating this inequality

$$x < f_1(x) < f_2(x) < \dots < f_j(x) < f(q) = q.$$

Clearly, the sequence converges and it must converge to a fixed point where $f(x) = x$. Similarly, if $q \leq x < 1$, then $f(q) \leq f(x) < x$ and iterating this inequality

$$x > f_1(x) > f_2(x) > \dots > f_j(x) > f(q) = q.$$

In the limit as j goes to infinity $f_j(x) = q$ for all x , $0 \leq x < 1$. ■

Recall that $f_j(x)$ is the generating function $\sum_{i=0}^{\infty} \text{Prob}(z_j = i) x^i$. The fact that in the limit the generating function equals the constant q , and is not a function of x , says that $\text{Prob}(z_j = 0) = q$ and $\text{Prob}(z_j = i) = 0$ for all finite nonzero values of i . The remaining probability is the probability of a nonfinite component. Thus, when $m > 1$, q is the extinction probability and $1-q$ is the probability that z_j grows without bound, i.e., immortality.

Theorem 3.11 Consider a tree generated by a branching process. Let $f(x)$ be the generating function for the number of children at each node.

1. If the expected number of children at each node is less than or equal to one, then the probability of extinction is one unless the probability of exactly one child is one.

2. If the expected number of children of each node is greater than one, then the probability of extinction is the unique solution to $f(x) = x$ in $[0, 1)$.

Proof: Let p_i be the probability of i children at each node. Then $f(x) = p_0 + p_1x + p_2x^2 + \dots$ is the generating function for the number of children at each node and $f'(1) = p_1 + 2p_2 + 3p_3 + \dots$ is the slope of $f(x)$ at $x = 1$. Observe that $f'(1)$ is the expected number of children at each node.

Since the expected number of children at each node is the slope of $f(x)$ at $x = 1$, if the expected number of children is less than or equal to one, the slope of $f(x)$ at $x = 1$ is less than or equal to one and the unique root of $f(x) = x$ in $(0, 1]$ is at $x = 1$ and the probability of extinction is one unless $f'(1) = 1$ and $p_1 = 1$. If $f'(1) = 1$ and $p_1 = 1$, $f(x) = x$ and the tree is an infinite degree one chain. If the slope of $f(x)$ at $x = 1$ is greater than one, then the probability of extinction is the unique solution to $f(x) = x$ in $[0, 1)$. ■

A branching process with $m < 1$ or $m = 1$ and $p_1 < 1$ dies out with probability one. If $m = 1$ and $p_1 = 1$, then the branching process consists of an infinite chain with no fan out. If $m > 1$, then the branching process will die out with some probability less than one unless $p_0 = 0$ in which case it cannot die out, since a node always has at least one descendent.

Note that the branching process corresponds to finding the size of a component in an infinite graph. In a finite graph, the probability distribution of descendants is not a constant as more and more vertices of the graph get discovered.

The simple branching process defined here either dies out or goes to infinity. In biological systems there are other factors, since processes often go to stable populations. One possibility is that the probability distribution for the number of descendants of a child depends on the total population of the current generation.

Expected size of extinct families

We now show that the expected size of an extinct family is finite, provided that $m \neq 1$. Note that at extinction, the size must be finite. However, the expected size at extinction could conceivably be infinite, if the probability of dying out did not decay fast enough. To see how the expected value of a random variable that is always finite could be infinite, let x be an integer valued random variable. Let p_i be the probability that $x = i$. If $\sum_{i=1}^{\infty} p_i = 1$, then with probability one, x will be finite. However, the expected value of x may be infinite. That is, $\sum_{i=0}^{\infty} ip_i = \infty$. For example, if for $i > 0$, $p_i = \frac{6}{\pi} \frac{1}{i^2}$, then $\sum_{i=1}^{\infty} p_i = 1$, but $\sum_{i=1}^{\infty} ip_i = \infty$. The value of the random variable x is always finite, but its expected value is infinite. This does not happen in a branching process, except in the special case where the slope $m = f'(1)$ equals one and $p_1 \neq 1$.

Lemma 3.12 *If the slope $m = f'(1)$ does not equal one, then the expected size of an extinct family is finite. If the slope m equals one and $p_1 = 1$, then the tree is an infinite degree one chain and there are no extinct families. If $m=1$ and $p_1 < 1$, then the expected size of the extinct family is infinite.*

Proof: Let z_i be the random variable denoting the size of the i^{th} generation and let q be the probability of extinction. The probability of extinction for a tree with k children in the first generation is q^k since each of the k children has an extinction probability of q . Note that the expected size of z_1 , the first generation, over extinct trees will be smaller than the expected size of z_1 over all trees since when the root node has a larger number of children than average, the tree is more likely to be infinite.

By Bayes rule

$$\text{Prob}(z_1 = k | \text{extinction}) = \text{Prob}(z_1 = k) \frac{\text{Prob}(\text{extinction} | z_1 = k)}{\text{Prob}(\text{extinction})} = p_k \frac{q^k}{q} = p_k q^{k-1}.$$

Knowing the probability distribution of z_1 given extinction, allows us to calculate the expected size of z_1 given extinction.

$$E(z_1 | \text{extinction}) = \sum_{k=0}^{\infty} k p_k q^{k-1} = f'(q).$$

We now prove, using independence, that the expected size of the i^{th} generation given extinction is

$$E(z_i | \text{extinction}) = \left(f'(q)\right)^i.$$

For $i = 2$, z_2 is the sum of z_1 independent random variables, each independent of the random variable z_1 . So, $E(z_2 | z_1 = j \text{ and extinction}) = E(\text{sum of } j \text{ copies of } z_1 | \text{extinction}) = jE(z_1 | \text{extinction})$. Summing over all values of j

$$\begin{aligned} E(z_2 | \text{extinction}) &= \sum_{j=1}^{\infty} E(z_2 | z_1 = j \text{ and extinction}) \text{Prob}(z_1 = j | \text{extinction}) \\ &= \sum_{j=1}^{\infty} j E(z_1 | \text{extinction}) \text{Prob}(z_1 = j | \text{extinction}) \\ &= E(z_1 | \text{extinction}) \sum_{j=1}^{\infty} j \text{Prob}(z_1 = j | \text{extinction}) = E^2(z_1 | \text{extinction}). \end{aligned}$$

Since $E(z_1 | \text{extinction}) = f'(q)$, $E(z_2 | \text{extinction}) = (f'(q))^2$. Similarly, $E(z_i | \text{extinction}) = (f'(q))^i$. The expected size of the tree is the sum of the expected sizes of each generation. That is,

$$\text{Expected size of tree given extinction} = \sum_{i=0}^{\infty} E(z_i | \text{extinction}) = \sum_{i=0}^{\infty} (f'(q))^i = \frac{1}{1 - f'(q)}.$$

Thus, the expected size of an extinct family is finite since $f'(q) < 1$ provided $m \neq 1$.

The fact that $f'(q) < 1$ is illustrated in Figure 3.10. If $m < 1$, then $q=1$ and $f'(q) = m$ is less than one. If $m > 1$, then $q \in [0, 1)$ and again $f'(q) < 1$ since q is the solution to $f(x) = x$ and $f'(q)$ must be less than one for the curve $f(x)$ to cross the line x . Thus, for $m < 1$ or $m > 1$, $f'(q) < 1$ and the expected tree size of $\frac{1}{1-f'(q)}$ is finite. For $m=1$ and $p_1 < 1$, one has $q=1$ and thus $f'(q) = 1$ and the formula for the expected size of the tree diverges. ■

3.5 Cycles and Full Connectivity

This section considers when cycles form and when the graph becomes fully connected. For both of these problems, we look at each subset of k vertices and see when they form either a cycle or a connected component.

3.5.1 Emergence of Cycles

The emergence of cycles in $G(n, p)$ has a threshold when p equals to $1/n$.

Theorem 3.13 *The threshold for the existence of cycles in $G(n, p)$ is $p = 1/n$.*

Proof: Let x be the number of cycles in $G(n, p)$. To form a cycle of length k , the vertices can be selected in $\binom{n}{k}$ ways. Given the k vertices of the cycle, they can be ordered by arbitrarily selecting a first vertex, then a second vertex in one of $k-1$ ways, a third in one of $k-2$ ways, etc. Since a cycle and its reversal are the same cycle, divide by 2. Thus, there are $\binom{n}{k} \frac{(k-1)!}{2}$ cycles of length k and

$$E(x) = \sum_{k=3}^n \binom{n}{k} \frac{(k-1)!}{2} p^k \leq \sum_{k=3}^n \frac{n^k}{2k} p^k \leq \sum_{k=3}^n (np)^k = (np)^3 \frac{1-(np)^{n-2}}{1-np} \leq 2(np)^3,$$

provided that $np < 1/2$. When p is asymptotically less than $1/n$, then $\lim_{n \rightarrow \infty} np = 0$ and

$\lim_{n \rightarrow \infty} \sum_{k=3}^n (np)^k = 0$. So, as n goes to infinity, $E(x)$ goes to zero. Thus, the graph almost surely has no cycles by the first moment method. A second moment argument can be used to show that for $p = d/n$, $d > 1$, a graph will have a cycle with probability tending to one. ■

The argument above does not yield a sharp threshold since we argued that $E(x) \rightarrow 0$ only under the assumption that p is asymptotically less than $\frac{1}{n}$. A sharp threshold requires $E(x) \rightarrow 0$ for $p = d/n$, $d < 1$.

Property	Threshold
cycles	$1/n$
giant component	$1/n$
giant component + isolated vertices	$\frac{1}{2} \frac{\ln n}{n}$
connectivity, disappearance of isolated vertices	$\frac{\ln n}{n}$
diameter two	$\sqrt{\frac{2 \ln n}{n}}$

Consider what happens in more detail when $p = d/n$, d a constant.

$$\begin{aligned}
E(x) &= \sum_{k=3}^n \binom{n}{k} \frac{(k-1)!}{2} p^k \\
&= \frac{1}{2} \sum_{k=3}^n \frac{n(n-1) \cdots (n-k+1)}{k!} (k-1)! p^k \\
&= \frac{1}{2} \sum_{k=3}^n \frac{n(n-1) \cdots (n-k+1)}{n^k} \frac{d^k}{k}.
\end{aligned}$$

$E(x)$ converges if $d < 1$, and diverges if $d \geq 1$. If $d < 1$, $E(x) \leq \frac{1}{2} \sum_{k=3}^n \frac{d^k}{k}$ and $\lim_{n \rightarrow \infty} E(x)$ equals a constant greater than zero. If $d = 1$, $E(x) = \frac{1}{2} \sum_{k=3}^n \frac{n(n-1) \cdots (n-k+1)}{n^k} \frac{1}{k}$. Consider only the first $\log n$ terms of the sum. Since $\frac{n}{n-i} = 1 + \frac{i}{n-i} \leq e^{i/n-i}$, it follows that $\frac{n(n-1) \cdots (n-k+1)}{n^k} \geq 1/2$. Thus,

$$E(x) \geq \frac{1}{2} \sum_{k=3}^{\log n} \frac{n(n-1) \cdots (n-k+1)}{n^k} \frac{1}{k} \geq \frac{1}{4} \sum_{k=3}^{\log n} \frac{1}{k}.$$

Then, in the limit as n goes to infinity

$$\lim_{n \rightarrow \infty} E(x) \geq \lim_{n \rightarrow \infty} \frac{1}{4} \sum_{k=3}^{\log n} \frac{1}{k} \geq \lim_{n \rightarrow \infty} (\log \log n) = \infty.$$

For $p = d/n$, $d < 1$, $E(x)$ converges to a nonzero constant and with some nonzero probability, graphs will have a constant number of cycles independent of the size of the graph. For $d > 1$, $E(x)$ converges to infinity and a second moment argument shows that graphs will have an unbounded number of cycles increasing with n .

3.5.2 Full Connectivity

As p increases from $p = 0$, small components form. At $p = 1/n$ a giant component emerges and swallows up smaller components, starting with the larger components and

ending up swallowing isolated vertices forming a single connected component at $p = \frac{\ln n}{n}$, at which point the graph becomes connected. We begin our development with a technical lemma.

Lemma 3.14 *The expected number of connected components of size k in $G(n, p)$ is at most*

$$\binom{n}{k} k^{k-2} p^{k-1} (1-p)^{kn-k^2}.$$

Proof: The probability that k vertices form a connected component consists of the product of two probabilities. The first is the probability that the k vertices are connected, and the second is the probability that there are no edges out of the component to the remainder of the graph. The first probability is at most the sum over all spanning trees of the k vertices, that the edges of the spanning tree are present. The "at most" in the lemma statement is because $G(n, p)$ may contain more than one spanning tree on these nodes and, in this case, the union bound is higher than the actual probability. There are k^{k-2} spanning trees on k nodes. See Section 11.8.7 in the appendix. The probability of all the $k-1$ edges of one spanning tree being present is p^{k-1} and the probability that there are no edges connecting the k vertices to the remainder of the graph is $(1-p)^{k(n-k)}$. Thus, the probability of one particular set of k vertices forming a connected component is at most $k^{k-2} p^{k-1} (1-p)^{kn-k^2}$. Thus, the expected number of connected components of size k is $\binom{n}{k} k^{k-2} p^{k-1} (1-p)^{kn-k^2}$. ■

We now prove that for $p = \frac{1}{2} \frac{\ln n}{n}$, the giant component has absorbed all small components except for isolated vertices.

Theorem 3.15 *Let $p = c \frac{\ln n}{n}$. For $c > 1/2$, almost surely there are only isolated vertices and a giant component. For $c > 1$, almost surely the graph is connected.*

Proof: We prove that almost surely for $c > 1/2$, there is no connected component with k vertices for any k , $2 \leq k \leq n/2$. This proves the first statement of the theorem since, if there were two or more components that are not isolated vertices, both of them could not be of size greater than $n/2$. The second statement that for $c > 1$ the graph is connected then follows from Theorem 3.6 which states that isolated vertices disappear at $c = 1$.

We now show that for $p = c \frac{\ln n}{n}$, the expected number of components of size k , $2 \leq k \leq n/2$, is less than n^{1-2c} and thus for $c > 1/2$ there are no components, except for isolated vertices and the giant component. Let x_k be the number of connected components of size k . Substitute $p = c \frac{\ln n}{n}$ into $\binom{n}{k} k^{k-2} p^{k-1} (1-p)^{kn-k^2}$ and simplify using $\binom{n}{k} \leq (en/k)^k$, $1-p \leq e^{-p}$, $k-1 < k$, and $x = e^{\ln x}$ to get

$$E(x_k) \leq \exp \left(\ln n + k + k \ln \ln n - 2 \ln k + k \ln c - ck \ln n + ck^2 \frac{\ln n}{n} \right).$$

Keep in mind that the leading terms here for large k are the last two and, in fact, at $k = n$, they cancel each other so that our argument does not prove the fallacious statement for $c \geq 1$ that there is no connected component of size n , since there is. Let

$$f(k) = \ln n + k + k \ln \ln n - 2 \ln k + k \ln c - ck \ln n + ck^2 \frac{\ln n}{n}.$$

Differentiating with respect to k ,

$$f'(k) = 1 + \ln \ln n - \frac{2}{k} + \ln c - c \ln n + \frac{2ck \ln n}{n}$$

and

$$f''(k) = \frac{2}{k^2} + \frac{2c \ln n}{n} > 0.$$

Thus, the function $f(k)$ attains its maximum over the range $[2, n/2]$ at one of the extreme points 2 or $n/2$. At $k = 2$, $f(2) \approx (1 - 2c) \ln n$ and at $k = n/2$, $f(n/2) \approx -c \frac{n}{4} \ln n$. So $f(k)$ is maximum at $k = 2$. For $k = 2$, $E(x)_k = e^{f(k)}$ is approximately $e^{(1-2c) \ln n} = n^{1-2c}$ and is geometrically falling as k increases from 2. At some point $E(x_k)$ starts to increase but never gets above $n^{-\frac{c}{4}n}$. Thus, the expected sum of the number of components of size k , for $2 \leq k \leq n/2$ is

$$E \left(\sum_{k=2}^{n/2} x_k \right) = O(n^{1-2c}).$$

This expected number goes to zero for $c > 1/2$ and the first-moment method implies that, almost surely, there are no components of size between 2 and $n/2$. This completes the proof of Theorem 3.15. ■

3.5.3 Threshold for $O(\ln n)$ Diameter

We now show that within a constant factor of the threshold for graph connectivity, not only is the graph connected, but its diameter is $O(\ln n)$. That is, if p is $\Omega(\ln n/n)$, the diameter of $G(n, p)$ is $O(\ln n)$.

Consider a particular vertex v . Let S_i be the set of vertices at distance i from v . We argue that as i grows, $|S_1| + |S_2| + \dots + |S_i|$ grows by a constant factor up to a size of $n/1000$. This implies that in $O(\ln n)$ steps, at least $n/1000$ vertices are connected to v . Then, there is a simple argument at the end of the proof of Theorem 3.17 that a pair of $n/1000$ sized subsets, connected to two different vertices v and w , have an edge between them.

Lemma 3.16 *Consider $G(n, p)$ for sufficiently large n with $p = c \ln n/n$ for any $c > 0$. Let S_i be the set of vertices at distance i from some fixed vertex v . If $|S_1| + |S_2| + \dots + |S_i| \leq n/1000$, then*

$$\text{Prob}(|S_{i+1}| < 2(|S_1| + |S_2| + \dots + |S_i|)) \leq e^{-10|S_i|}.$$

Proof: Let $|S_i| = k$. For each vertex u not in $S_1 \cup S_2 \cup \dots \cup S_i$, the probability that u is not in S_{i+1} is $(1-p)^k$ and these events are independent. So, $|S_{i+1}|$ is the sum of $n - (|S_1| + |S_2| + \dots + |S_i|)$ independent Bernoulli random variables, each with probability of

$$1 - (1-p)^k \geq 1 - e^{-ck \ln n/n}$$

of being one. Note that $n - (|S_1| + |S_2| + \dots + |S_i|) \geq 999n/1000$. So,

$$E(|S_{i+1}|) \geq \frac{999n}{1000} (1 - e^{-ck \frac{\ln n}{n}}).$$

Subtracting $200k$ from each side

$$E(|S_{i+1}|) - 200k \geq \frac{n}{2} \left(1 - e^{-ck \frac{\ln n}{n}} - 400 \frac{k}{n} \right).$$

Let $\alpha = \frac{k}{n}$ and $f(\alpha) = 1 - e^{-c\alpha \ln n} - 400\alpha$. By differentiation $f''(\alpha) \leq 0$, so f is concave and the minimum value of f over the interval $[0, 1/1000]$ is attained at one of the end points. It is easy to check that both $f(0)$ and $f(1/1000)$ are greater than or equal to zero for sufficiently large n . Thus, f is nonnegative throughout the interval proving that $E(|S_{i+1}|) \geq 200|S_i|$. The lemma follows from Chernoff bounds. ■

Theorem 3.17 For $p \geq c \ln n/n$, where c is a sufficiently large constant, almost surely, $G(n, p)$ has diameter $O(\ln n)$.

Proof: By Corollary 3.2, almost surely, the degree of every vertex is $\Omega(np) = \Omega(\ln n)$, which is at least $20 \ln n$ for c sufficiently large. Assume this holds. So, for a fixed vertex v , S_1 as defined in Lemma 3.16 satisfies $|S_1| \geq 20 \ln n$.

Let i_0 be the least i such that $|S_1| + |S_2| + \dots + |S_i| > n/1000$. From Lemma 3.16 and the union bound, the probability that for some i , $1 \leq i \leq i_0 - 1$, $|S_{i+1}| < 2(|S_1| + |S_2| + \dots + |S_i|)$ is at most $\sum_{k=20 \ln n}^{n/1000} e^{-10k} \leq 1/n^4$. So, with probability at least $1 - (1/n^4)$, each S_{i+1} is at least double the sum of the previous S_j 's, which implies that in $O(\ln n)$ steps, $i_0 + 1$ is reached.

Consider any other vertex w . We wish to find a short $O(\ln n)$ length path between v and w . By the same argument as above, the number of vertices at distance $O(\ln n)$ from w is at least $n/1000$. To complete the argument, either these two sets intersect in which case we have found a path from v to w of length $O(\ln n)$ or they do not intersect. In the latter case, with high probability there is some edge between them. For a pair of disjoint sets of size at least $n/1000$, the probability that none of the possible $n^2/10^6$ or more edges between them is present is at most $(1-p)^{n^2/10^6} = e^{-\Omega(n \ln n)}$. There are at most 2^{2n} pairs of such sets and so the probability that there is some such pair with no edges is $e^{-\Omega(n \ln n) + O(n)} \rightarrow 0$. Note that there is no conditioning problem since we are arguing this for every pair of such sets. Think of whether such an argument made for just the n subsets of vertices, which are vertices at distance at most $O(\ln n)$ from a specific vertex, would work. ■

3.6 Phase Transitions for Increasing Properties

For many graph properties such as connectivity, having no isolated vertices, having a cycle, etc., the probability of a graph having the property increases as edges are added to the graph. Such a property is called an increasing property. Q is an *increasing property* of graphs if when a graph G has the property, any graph obtained by adding edges to G must also have the property. In this section we show that any increasing property, in fact, has a threshold, although not necessarily a sharp one.

The notion of increasing property is defined in terms of adding edges. The following lemma proves that if Q is an increasing property, then increasing p in $G(n, p)$ increases the probability of the property Q .

Lemma 3.18 *If Q is an increasing property of graphs and $0 \leq p \leq q \leq 1$, then the probability that $G(n, q)$ has property Q is greater than or equal to the probability that $G(n, p)$ has property Q .*

Proof: This proof uses an interesting relationship between $G(n, p)$ and $G(n, q)$. Generate $G(n, q)$ as follows. First generate $G(n, p)$. This means generating a graph on n vertices with edge probabilities p . Then, independently generate another graph $G\left(n, \frac{q-p}{1-p}\right)$ and take the union by putting in an edge if either of the two graphs has the edge. Call the resulting graph H . The graph H has the same distribution as $G(n, q)$. This follows since the probability that an edge is in H is $p + (1-p)\frac{q-p}{1-p} = q$, and, clearly, the edges of H are independent. The lemma follows since whenever $G(n, p)$ has the property Q , H also has the property Q . ■

We now introduce a notion called *replication*. An m -fold replication of $G(n, p)$ is a random graph obtained as follows. Generate m independent copies of $G(n, p)$. Include an edge in the m -fold replication if the edge is in any one of the m copies of $G(n, p)$. The resulting random graph has the same distribution as $G(n, q)$ where $q = 1 - (1-p)^m$ since the probability that a particular edge is not in the m -fold replication is the product of probabilities that it is not in any of the m copies of $G(n, p)$. If the m -fold replication of $G(n, p)$ does not have an increasing property Q , then none of the m copies of $G(n, p)$ has the property. The converse is not true. If no copy has the property, their union may have it. Since Q is an increasing property and $q = 1 - (1-p)^m \leq 1 - (1-mp) = mp$

$$\text{Prob}(G(n, mp) \text{ has } Q) \geq \text{Prob}(G(n, q) \text{ has } Q) \quad (3.3)$$

We now show that any increasing property Q has a phase transition. The transition occurs at the point at which the probability that $G(n, p)$ has property Q is $\frac{1}{2}$. We will prove that for any function asymptotically less than $p(n)$ that the probability of having property Q goes to zero as n goes to infinity.

Theorem 3.19 *Any increasing property Q of $G(n, p)$ has a phase transition at $p(n)$, where for each n , $p(n)$ is the minimum real number a_n for which the probability that $G(n, a_n)$ has property Q is $1/2$.*

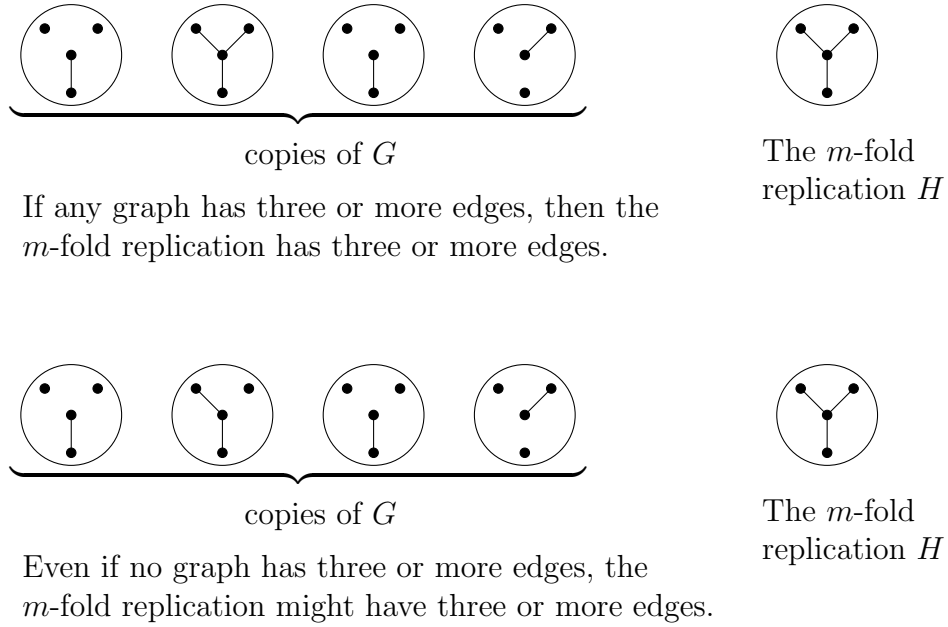


Figure 3.12: The property that G has three or more edges is an increasing property. Let H be the m -fold replication of G . If any copy of G has three or more edges, H has three or more edges. However, H can have three or more edges even if no copy of G has three or more edges.

Proof: Let $p_0(n)$ be any function such that

$$\lim_{n \rightarrow \infty} \frac{p_0(n)}{p(n)} = 0.$$

We assert that almost surely $G(n, p_0)$ does not have the property Q . Suppose for contradiction, that this is not true. That is, the probability that $G(n, p_0)$ has the property Q does not converge to zero. By the definition of a limit, there exists $\varepsilon > 0$ for which the probability that $G(n, p_0)$ has property Q is at least ε on an infinite set I of n . Let $m = \lceil (1/\varepsilon) \rceil$. Let $G(n, q)$ be the m -fold replication of $G(n, p_0)$. The probability that $G(n, q)$ does not have Q is at most $(1 - \varepsilon)^m \leq e^{-1} \leq 1/2$ for all $n \in I$. For these n , by (3.3)

$$\text{Prob}(G(n, mp_0) \text{ has } Q) \geq \text{Prob}(G(n, q) \text{ has } Q) \geq 1/2.$$

Since $p(n)$ is the minimum real number a_n for which the probability that $G(n, a_n)$ has property Q is $1/2$, it must be that $mp_0(n) \geq p(n)$. This implies that $\frac{p_0(n)}{p(n)}$ is at least $1/m$ infinitely often, contradicting the hypothesis that $\lim_{n \rightarrow \infty} \frac{p_0(n)}{p(n)} = 0$.

A symmetric argument shows that for any $p_1(n)$ such that $\lim_{n \rightarrow \infty} \frac{p(n)}{p_1(n)} = 0$, $G(n, p_1)$ almost surely has property Q . ■

3.7 Phase Transitions for CNF-sat

Phase transitions occur not only in random graphs, but in other random structures as well. An important example is that of satisfiability for a Boolean formula in conjunctive normal form.

Generate a random CNF formula f with n variables, m clauses, and k literals per clause. Each clause is picked independently with k literals picked uniformly at random from the set of $2n$ possible literals to form the clause. Here, the number of clauses n is going to infinity, m is a function of n , and k is a fixed constant. A reasonable value to think of for k is $k = 3$. A literal is a variable or its negation. Unsatisfiability is an increasing property since adding more clauses preserves unsatisfiability. By arguments similar to the last section, there is a phase transition, i.e., a function $m(n)$ such that if $m_1(n)$ is $o(m(n))$, a random formula with $m_1(n)$ clauses is, almost surely, satisfiable and for $m_2(n)$ with $m_2(n)/m(n) \rightarrow \infty$, a random formula with $m_2(n)$ clauses is, almost surely, unsatisfiable. It has been conjectured that there is a constant r_k independent of n such that $r_k n$ is a sharp threshold.

Here we derive upper and lower bounds on r_k . It is relatively easy to get an upper bound on r_k . A fixed truth assignment satisfies a random k clause with probability $1 - \frac{1}{2^k}$. Of the 2^k truth assignments to the k variables in the clause, only one fails to satisfy the clause. Thus, with probability $\frac{1}{2^k}$, the clause is not satisfied, and with probability $1 - \frac{1}{2^k}$, the clause is satisfied. Let $m = cn$. Now, cn independent clauses are all satisfied by the fixed assignment with probability $(1 - \frac{1}{2^k})^{cn}$. Since there are 2^n truth assignments, the expected number of satisfying assignments for a formula with cn clauses is $2^n (1 - \frac{1}{2^k})^{cn}$. If $c = 2^k \ln 2$, the expected number of satisfying assignments is

$$2^n \left(1 - \frac{1}{2^k}\right)^{n2^k \ln 2}.$$

$(1 - \frac{1}{2^k})^{2^k}$ is at most $1/e$ and approaches $1/e$ in the limit. Thus,

$$2^n \left(1 - \frac{1}{2^k}\right)^{n2^k \ln 2} \leq 2^n e^{-n \ln 2} = 2^n 2^{-n} = 1.$$

For $c > 2^k \ln 2$, the expected number of satisfying assignments goes to zero as $n \rightarrow \infty$. Here the expectation is over the choice of clauses which is random, not the choice of a truth assignment. From the first moment method, it follows that a random formula with cn clauses is almost surely not satisfiable. Thus, $r_k \leq 2^k \ln 2$.

The other direction, showing a lower bound for r_k , is not that easy. From now on, we focus only on the case $k = 3$. The statements and algorithms given here can be extended to $k \geq 4$, but with different constants. It turns out that the second moment method cannot be directly applied to get a lower bound on r_3 because the variance is too high. A simple algorithm, called the Smallest Clause Heuristic (abbreviated SC), yields a satisfying assignment with probability tending to one if $c < \frac{2}{3}$, proving that $r_3 \geq \frac{2}{3}$. Other more

difficult to analyze algorithms, push the lower bound on r_3 higher.

The Smallest Clause Heuristic repeatedly executes the following. Assign true to a random literal in a random smallest length clause and delete the clause since it is now satisfied. Pick at random a 1-literal clause, if one exists, and set that literal to true. If there is no 1-literal clause, pick a 2-literal clause, select one of its two literals and set the literal to true. Otherwise, pick a 3-literal clause and a literal in it and set the literal to true. If we encounter a 0-length clause, then we have failed to find a satisfying assignment; otherwise, we have found one.

A related heuristic, called the Unit Clause Heuristic, selects a random clause with one literal, if there is one, and sets the literal in it to true. Otherwise, it picks a random as yet unset literal and sets it to true. The “pure literal” heuristic sets a random “pure literal”, a literal whose negation does not occur in any clause, to true, if there are any pure literals; otherwise, it sets a random literal to true.

When a literal w is set to true, all clauses containing w are deleted, since they are satisfied, and $\bar{w}$ is deleted from any clause containing $\bar{w}$. If a clause is reduced to length zero (no literals), then the algorithm has failed to find a satisfying assignment to the formula. The formula may, in fact, be satisfiable, but the algorithm has failed.

Example: Consider a 3-CNF formula with n variables and cn clauses. With n variables there are $2n$ literals, since a variable and its complement are distinct literals. The expected number of times a literal occurs is calculated as follows. Each clause has three literals. Thus, each of the $2n$ different literals occurs $\frac{(3cn)}{2n} = \frac{3}{2}c$ times on average. Suppose $c = 5$. Then each literal appears 7.5 times on average. If one sets a literal to true, one would expect to satisfy 7.5 clauses. However, this process is not repeatable since after setting a literal to true there is conditioning so that the formula is no longer random. ■

Theorem 3.20 *If the number of clauses in a random 3-CNF formula grows as cn where c is a constant less than $2/3$, then with probability $1 - o(1)$, the Shortest Clause Heuristic finds a satisfying assignment.*

The proof of this theorem will take the rest of the section. A general impediment to proving that simple algorithms work for random instances of many problems is conditioning. At the start, the input is random and has properties enjoyed by random instances. But, as the algorithm is executed; the data is no longer random, it is conditioned on the steps of the algorithm so far. In the case of SC and other heuristics for finding a satisfying assignment for a Boolean formula, the argument to deal with conditioning is relatively simple.

We supply some intuition before going to the proof. Imagine maintaining a queue of 1 and 2-clauses. A 3-clause enters the queue when one of its literals is set to false and it becomes a 2-clause. SC always picks a 1 or 2-clause if there is one and sets one of its

literals to true. At any step when the total number of 1 and 2-clauses is positive, one of the clauses is removed from the queue. Consider the arrival rate, the expected number of arrivals into the queue. For a particular clause to arrive into the queue at time t to become a 2-clause, it must contain the negation of the literal being set to true at time t . It can contain any two other literals not yet set. The number of such clauses is $\binom{n-t}{2}2^2$. So, the probability that a particular clause arrives in the queue at time t is at most

$$\frac{\binom{n-t}{2}2^2}{\binom{n}{3}2^3} \leq \frac{3}{2(n-2)}.$$

Since there are cn clauses in total, the arrival rate is $\frac{3c}{2}$, which for $c < 2/3$ is a constant strictly less than one. The arrivals into the queue of different clauses occur independently (Lemma 3.21), the queue has arrival rate strictly less than one, and the queue loses one or more clauses whenever it is nonempty. This implies that the queue never has too many clauses in it. A slightly more complicated argument will show that no clause remains as a 1 or 2-clause for $\Omega(\ln n)$ steps (Lemma 3.22). This implies that the probability of two contradictory 1-length clauses, which is a precursor to a 0-length clause, is very small.

Lemma 3.21 *Let T_i be the first time that clause i turns into a 2-clause. T_i is ∞ if clause i gets satisfied before turning into a 2-clause. The T_i are mutually independent and for any t ,*

$$\text{Prob}(T_i = t) \leq \frac{3}{2(n-2)}.$$

Proof: For the proof, generate the clauses in a different way. The important thing is that the new method of generation, called the method of “deferred decisions”, results in the same distribution of input formulae as the original. The method of deferred decisions is tied in with the SC algorithm and works as follows. At any time, the length of each clause (number of literals) is all that we know; we have not yet picked which literals are in each clause. At the start, every clause has length three and SC picks one of the clauses uniformly at random. Now, SC wants to pick one of the three literals in that clause to set to true, but we do not know which literals are in the clause. At this point, we pick uniformly at random one of the $2n$ possible literals. Say for illustration, we picked $\bar{x}_{102}$. The literal $\bar{x}_{102}$ is placed in the clause and set to true. The literal x_{102} is set to false. We must also deal with occurrences of the literal or its negation in all other clauses, but again, we do not know which clauses have such an occurrence. We decide that now. For each clause, independently, with probability $3/n$, include the variable x_{102} or $\bar{x}_{102}$ in the clause and if included, with probability $1/2$, include the literal $\bar{x}_{102}$ in the clause and with the other $1/2$ probability include its negation, namely, x_{102} . In either case, we decrease the residual length of the clause by one. The algorithm deletes the clause since it is satisfied and we do not care which other literals are in it. If we had included the negation of the literal instead, then we delete just that occurrence, and decrease the length of the clause by one.

At a general stage, suppose the fates of i variables have already been decided and $n - i$ remain. The residual length of each clause is known. Among the clauses that are not yet satisfied, choose a random shortest length clause. Among the $n - i$ variables remaining, pick one uniformly at random, then pick it or its negation as the new literal. Include this literal in the clause thereby satisfying it. Since the clause is satisfied, the algorithm deletes it. For each other clause, do the following. If its residual length is l , decide with probability $l/(n - i)$ to include the new variable in the clause and if so with probability $1/2$ each, include it or its negation. If literal v is included in a clause, delete the clause as it is now satisfied. If $\bar{v}$ is included in a clause, then just delete the literal and decrease the residual length of the clause by one.

Why does this yield the same distribution as the original one? First, observe that the order in which the variables are picked by the method of deferred decisions is independent of the clauses; it is just a random permutation of the n variables. Look at any one clause. For a clause, we decide in order whether each variable or its negation is in the clause. So for a particular clause and a particular triple i, j , and k with $i < j < k$, the probability that the clause contains the i^{th} , the j^{th} , and k^{th} literal (or their negations) in the order determined by deferred decisions is:

$$\begin{aligned} & \left(1 - \frac{3}{n}\right) \left(1 - \frac{3}{n-1}\right) \cdots \left(1 - \frac{3}{n-i+2}\right) \frac{3}{n-i+1} \\ & \left(1 - \frac{2}{n-i}\right) \left(1 - \frac{2}{n-i-1}\right) \cdots \left(1 - \frac{2}{n-j+2}\right) \frac{2}{n-j+1} \\ & \left(1 - \frac{1}{n-j}\right) \left(1 - \frac{1}{n-j-1}\right) \cdots \left(1 - \frac{1}{n-k+2}\right) \frac{1}{n-k+1} = \frac{3}{n(n-1)(n-2)}, \end{aligned}$$

where the $(1 - \cdots)$ factors are for not picking the current variable or negation to be included and the others are for including the current variable or its negation. Independence among clauses follows from the fact that we have never let the occurrence or nonoccurrence of any variable in any clause influence our decisions on other clauses.

Now, we prove the lemma by appealing to the method of deferred decisions to generate the formula. $T_i = t$ if and only if the method of deferred decisions does not put the current literal at steps $1, 2, \dots, t - 1$ into the i^{th} clause, but puts the negation of the literal at step t into it. Thus, the probability is precisely

$$\frac{1}{2} \left(1 - \frac{3}{n}\right) \left(1 - \frac{3}{n-1}\right) \cdots \left(1 - \frac{3}{n-t+2}\right) \frac{3}{n-t+1} \leq \frac{3}{2(n-2)},$$

as claimed. Clearly the T_i are independent since again deferred decisions deal with different clauses independently. ■

Lemma 3.22 *With probability $1 - o(1)$, no clause remains a 2 or 1-clause for $\Omega(\ln n)$ steps. I.e., once a 3-clause becomes a 2-clause, it is either satisfied or reduced to a 0-clause in $O(\ln n)$ steps.*

Proof: Without loss of generality, again focus on the first clause. Suppose it becomes a 2-clause at step s_1 and remains a 2 or 1-clause until step s . Suppose $s - s_1 \geq c_2 \ln n$. Let r

be the last time before s when there are no 2 or 1-clauses at all. Since at time 0, there are no 2 or 1-clauses, r is well-defined. We have $s - r \geq c_2 \ln n$. In the interval r to s , at each step, there is at least one 2 or 1-clause. Since SC always decreases the total number of 1 and 2-clauses by one whenever it is positive, we must have generated at least $s - r$ new 2-clauses between r and s . Now, define an indicator random variable for each 3-clause which has value one if the clause turns into a 2-clause between r and s . By Lemma 3.21 these variables are independent and the probability that a particular 3-clause turns into a 2-clause at a time t is at most $3/(2(n - 2))$. Summing over t between r and s ,

$$\text{Prob (a 3-clause turns into a 2-clause during } [r, s]) \leq \frac{3(s - r)}{2(n - 2)}.$$

Since there are cn clauses in all, the expected sum of the indicator random variables is $cn \frac{3(s-r)}{2(n-2)} \approx \frac{3c(s-r)}{2}$. Note that $3c/2 < 1$, which implies the arrival rate into the queue of 2 and 1-clauses is a constant strictly less than one. Using Chernoff bounds, the probability that more than $s - r$ clauses turn into 2-clauses between r and s is at most $o(1/n^5)$. This is for one choice of a clause, one choice of s_1 and one choice each of r and s within $O(\ln n)$ of s_1 . Applying the union bound over $O(n^3)$ choices of clauses, $O(n)$ choices of s_1 and $O(\ln n)^2$ choices of r and s , we get that the probability that any clause remains a 2 or 1-clause for $\Omega(\ln n)$ steps is $o(1)$. ■

Now, suppose SC terminates in failure. At some time t , the algorithm generates a 0-clause. At time $t - 1$, this clause must have been a 1-clause. Suppose the clause consists of the literal w . Since at time $t - 1$, there is at least one 1-clause, the shortest clause rule of SC selects a 1-clause and sets the literal in that clause to true. This other clause must have been $\bar{w}$. Let t_1 be the first time either of these two clauses, w or $\bar{w}$, became a 2-clause. We have $t - t_1 \in O(\ln n)$. Clearly, until time t , neither of these two clauses is picked by SC. So, the literals which are set to true during this period are chosen independent of these clauses. Say the two clauses were $w + x + y$ and $\bar{w} + u + v$ at the start. x, y, u , and v must all be negations of literals set to true during steps t_1 to t . So, there are only $O((\ln n)^4)$ choices for x, y, u , and v . There are $O(n)$ choices of w , $O(n^2)$ choices of which two clauses of the input become these w and $\bar{w}$, and n choices for t_1 . Thus, there are $O(n^4(\ln n)^4)$ choices for these clauses. The probability of these choices is therefore $O(n^4(\ln n)^4/n^6) = o(1)$, as required.

3.8 Nonuniform and Growth Models of Random Graphs

3.8.1 Nonuniform Models

So far we have considered the random graph $G(n, p)$ in which all vertices have the same expected degree and showed that the degree is concentrated close to its expectation. However, large graphs occurring in the real world tend to have power law degree distributions. For a power law degree distribution, the number $f(d)$ of vertices of degree d plotted as a function of d satisfies $f(d) \leq c/d^\alpha$, where α and c are constants.

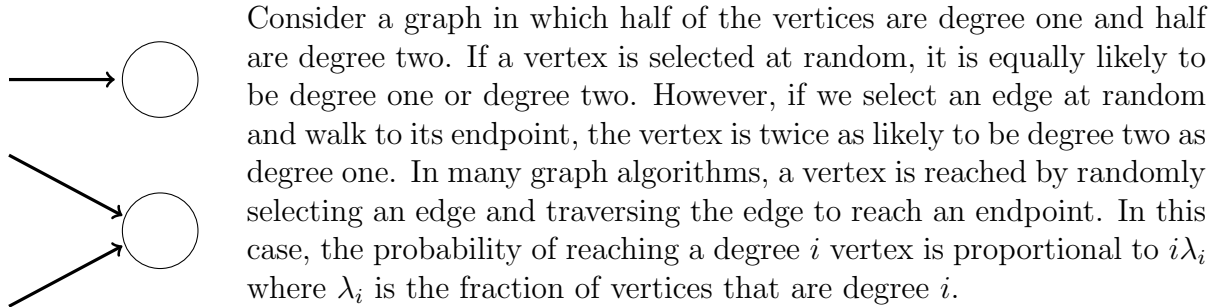


Figure 3.13: Probability of encountering a degree d vertex when following a path in a graph.

To generate such graphs, we stipulate that there are $f(d)$ vertices of degree d and choose uniformly at random from the set of graphs with this degree distribution. Clearly, in this model the graph edges are not independent and this makes these random graphs harder to analyze. But the question of when phase transitions occur in random graphs with arbitrary degree distributions is still of interest. In this section, we consider when a random graph with a nonuniform degree distribution has a giant component. Our treatment in this section, and subsequent ones, will be more intuitive without providing rigorous proofs.

3.8.2 Giant Component in Random Graphs with Given Degree Distribution

Molloy and Reed address the issue of when a random graph with a nonuniform degree distribution has a giant component. Let λ_i be the fraction of vertices of degree i . There will be a giant component if and only if $\sum_{i=0}^{\infty} i(i-2)\lambda_i > 0$.

To see intuitively that this is the correct formula, consider exploring a component of a graph starting from a given seed vertex. Degree zero vertices do not occur except in the case where the vertex is the seed. If a degree one vertex is encountered, then that terminates the expansion along the edge into the vertex. Thus, we do not want to encounter too many degree one vertices. A degree two vertex is neutral in that the vertex is entered by one edge and left by the other. There is no net increase in the size of the frontier. Vertices of degree i greater than two increase the frontier by $i-2$ vertices. The vertex is entered by one of its edges and thus there are $i-1$ edges to new vertices in the frontier for a net gain of $i-2$. The $i\lambda_i$ in $i(i-2)\lambda_i$ is proportional to the probability of reaching a degree i vertex and the $i-2$ accounts for the increase or decrease in size of the frontier when a degree i vertex is reached.

Example: Consider applying the Molloy Reed conditions to the $G(n, p)$ model. The summation $\sum_{i=0}^n i(i-2)p_i$ gives value zero precisely when $p = 1/n$, the point at which the phase transition occurs. At $p = 1/n$, the average degree of each vertex is one and

there are $n/2$ edges. However, the actual degree distribution of the vertices is binomial, where the probability that a vertex is of degree i is given by $p_i = \binom{n}{i} p^i (1-p)^{n-i}$. We now show that $\lim_{n \rightarrow \infty} \sum_{i=0}^n i(i-2)p_i = 0$ for $p_i = \binom{n}{i} p^i (1-p)^{n-i}$ when $p = 1/n$.

$$\begin{aligned}
& \lim_{n \rightarrow \infty} \sum_{i=0}^n i(i-2) \binom{n}{i} \left(\frac{1}{n}\right)^i \left(1 - \frac{1}{n}\right)^{n-i} \\
&= \lim_{n \rightarrow \infty} \sum_{i=0}^n i(i-2) \frac{n(n-1) \cdots (n-i+1)}{i! n^i} \left(1 - \frac{1}{n}\right)^n \left(1 - \frac{1}{n}\right)^{-i} \\
&= \frac{1}{e} \lim_{n \rightarrow \infty} \sum_{i=0}^n i(i-2) \frac{n(n-1) \cdots (n-i+1)}{i! n^i} \left(\frac{n}{n-1}\right)^i \\
&\leq \sum_{i=0}^{\infty} \frac{i(i-2)}{i!}.
\end{aligned}$$

To see that $\sum_{i=0}^{\infty} \frac{i(i-2)}{i!} = 0$, note that

$$\sum_{i=0}^{\infty} \frac{i}{i!} = \sum_{i=1}^{\infty} \frac{i}{i!} = \sum_{i=1}^{\infty} \frac{1}{(i-1)!} = \sum_{i=0}^{\infty} \frac{1}{i!}$$

and

$$\sum_{i=0}^{\infty} \frac{i^2}{i!} = \sum_{i=1}^{\infty} \frac{i}{(i-1)!} = \sum_{i=0}^{\infty} \frac{i+1}{i!} = \sum_{i=0}^{\infty} \frac{i}{i!} + \sum_{i=0}^{\infty} \frac{1}{i!} = 2 \sum_{i=0}^{\infty} \frac{1}{i!}.$$

Thus,

$$\sum_{i=0}^{\infty} \frac{i(i-2)}{i!} = \sum_{i=0}^{\infty} \frac{i^2}{i!} - 2 \sum_{i=0}^{\infty} \frac{i}{i!} = 0.$$

■

3.9 Growth Models

3.9.1 Growth Model Without Preferential Attachment

Many graphs that arise in the outside world started as small graphs that grew over time. In a model for such graphs, vertices and edges are added to the graph over time. In such a model there are many ways in which to select the vertices for attaching a new edge. One is to select two vertices uniformly at random from the set of existing vertices. Another is to select two vertices with probability proportional to their degree. This latter method is referred to as preferential attachment. A variant of this method

would be to add a new vertex at each unit of time and with probability δ add an edge where one end of the edge is the new vertex and the other end is a vertex selected with probability proportional to its degree. The graph generated by this latter method is a tree.

Consider a growth model for a random graph without preferential attachment. Start with zero vertices at time zero. At each unit of time a new vertex is created and with probability δ , two vertices chosen at random are joined by an edge. The two vertices may already have an edge between them. In this case, we will add another edge. So, the resulting structure is a multi-graph, rather than a graph. But since at time t , there are t vertices and in expectation only $O(\delta t)$ edges where there are t^2 pairs of vertices, it is very unlikely that there will be multiple edges.

The degree distribution for this growth model is calculated as follows. The number of vertices of degree k at time t is a random variable. Let $d_k(t)$ be the expectation of the number of vertices of degree k at time t . The number of isolated vertices increases by one at each unit of time and decreases by the number of isolated vertices, $b(t)$, that are picked to be end points of the new edge. $b(t)$ can take on values 0, 1, or 2. Taking expectations,

$$d_0(t+1) = d_0(t) + 1 - E(b(t)).$$

Now $b(t)$ is the sum of two 0-1 valued random variables whose values are the number of degree zero vertices picked for each end point of the new edge. Even though the two random variables are not independent, the expectation of $b(t)$ is the sum of the expectations of the two variables and is $2\delta \frac{d_0(t)}{t}$. Thus,

$$d_0(t+1) = d_0(t) + 1 - 2\delta \frac{d_0(t)}{t}.$$

The number of degree k vertices increases whenever a new edge is added to a degree $k-1$ vertex and decreases when a new edge is added to a degree k vertex. Reasoning as above,

$$d_k(t+1) = d_k(t) + 2\delta \frac{d_{k-1}(t)}{t} - 2\delta \frac{d_k(t)}{t}. \quad (3.4)$$

Note that this formula, as others in this section, is not quite precise. For example, the same vertex may be picked twice, so that the new edge is a self-loop. For $k \ll t$, this problem contributes a minuscule error. Restricting k to be a fixed constant and letting $t \rightarrow \infty$ in this section avoids these problems.

Assume that the above equations are exactly valid. Clearly, $d_0(1) = 1$ and $d_1(1) = d_2(1) = \dots = 0$. By induction on t , there is a unique solution to (3.4), since given $d_k(t)$ for all k , the equation determines $d_k(t+1)$ for all k . There is a solution of the form $d_k(t) = p_k t$, where p_k depends only on k and not on t , provided k is fixed and $t \rightarrow \infty$. Again, this is not precisely true, $d_1(1) = 0$ and $d_1(2) > 0$ clearly contradict the existence of a solution of the form $d_1(t) = p_1 t$.

Set $d_k(t) = p_k t$. Then,

$$(t+1)p_0 = p_0 t + 1 - 2\delta \frac{p_0 t}{t}$$

$$p_0 = 1 - 2\delta p_0$$

$$p_0 = \frac{1}{1+2\delta}$$

and

$$(t+1)p_k = p_k t + 2\delta \frac{p_{k-1} t}{t} - 2\delta \frac{p_k t}{t}$$

$$p_k = 2\delta p_{k-1} - 2\delta p_k$$

$$\begin{aligned} p_k &= \frac{2\delta}{1+2\delta} p_{k-1} \\ &= \left(\frac{2\delta}{1+2\delta} \right)^k p_0 \\ &= \frac{1}{1+2\delta} \left(\frac{2\delta}{1+2\delta} \right)^k. \end{aligned} \tag{3.5}$$

Thus, the model gives rise to a graph with a degree distribution that falls off exponentially fast with degree.

The generating function for component size

Let $n_k(t)$ be the expected number of components of size k at time t . Then $n_k(t)$ is proportional to the probability that a randomly picked component is of size k . This is not the same as picking the component containing a randomly selected vertex (see Figure 3.14). Indeed, the probability that the size of the component containing a randomly selected vertex is k is proportional to $kn_k(t)$. We will show that there is a solution for $n_k(t)$ of the form $a_k t$ where a_k is a constant independent of t . After showing this, we focus on the generating function $g(x)$ for the numbers $ka_k(t)$ and use $g(x)$ to find the threshold for giant components.

Consider $n_1(t)$, the expected number of isolated vertices at time t . At each unit of time, an isolated vertex is added to the graph and an expected $\frac{2\delta n_1(t)}{t}$ many isolated vertices are chosen for attachment and thereby leave the set of isolated vertices. Thus,

$$n_1(t+1) = n_1(t) + 1 - 2\delta \frac{n_1(t)}{t}.$$

For $k > 1$, $n_k(t)$ increases when two smaller components whose sizes sum to k are joined by an edge and decreases when a vertex in a component of size k is chosen for attachment.

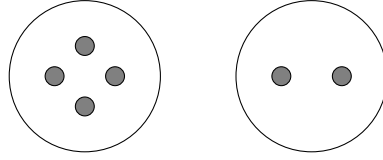


Figure 3.14: In selecting a component at random, each of the two components is equally likely to be selected. In selecting the component containing a random vertex, the larger component is twice as likely to be selected.

The probability that a vertex selected at random will be in a size k component is $\frac{kn_k(t)}{t}$. Thus,

$$n_k(t+1) = n_k(t) + \delta \sum_{j=1}^{k-1} \frac{jn_j(t)}{t} \frac{(k-j)n_{k-j}(t)}{t} - 2\delta \frac{kn_k(t)}{t}.$$

To be precise, one needs to consider the actual number of components of various sizes, rather than the expected numbers. Also, if both vertices at the end of the edge are in the same k -vertex component, then $n_k(t)$ does not go down as claimed. These small inaccuracies can be ignored.

Consider solutions of the form $n_k(t) = a_k t$. Note that $n_k(t) = a_k t$ implies the number of vertices in a connected component of size k is $ka_k t$. Since the total number of vertices at time t is t , ka_k is the probability that a random vertex is in a connected component of size k . The recurrences here are valid only for k fixed as $t \rightarrow \infty$. So $\sum_{k=0}^{\infty} ka_k$ may be less than 1, in which case, there are nonfinite size components whose sizes are growing with t . Solving for a_k yields $a_1 = \frac{1}{1+2\delta}$ and $a_k = \frac{\delta}{1+2k\delta} \sum_{j=1}^{k-1} j(k-j)a_j a_{k-j}$.

Consider the generating function $g(x)$ for the distribution of component sizes where the coefficient of x^k is the probability that a vertex chosen at random is in a component of size k .

$$g(x) = \sum_{k=1}^{\infty} ka_k x^k.$$

Now, $g(1) = \sum_{k=0}^{\infty} ka_k$ is the probability that a randomly chosen vertex is in a finite sized component. For $\delta = 0$, this is clearly one, since all vertices are in components of size one. On the other hand, for $\delta = 1$, the vertex created at time one has expected degree $\log n$, so it is in a nonfinite size component. This implies that for $\delta = 1$, $g(1) < 1$ and there is a nonfinite size component. Assuming continuity, there is a $\delta_{critical}$ above which $g(1) < 1$. From the formula for the a'_i s, we will derive the differential equation

$$g = -2\delta x g' + 2\delta x g g' + x$$

and then use the equation for g to determine the value of δ at which the phase transition for the appearance of a nonfinite sized component occurs.

Derivation of $g(x)$

From

$$a_1 = \frac{1}{1 + 2\delta}$$

and

$$a_k = \frac{\delta}{1 + 2k\delta} \sum_{j=1}^{k-1} j(k-j)a_j a_{k-j}$$

derive the equations

$$a_1 (1 + 2\delta) - 1 = 0$$

and

$$a_k (1 + 2k\delta) = \delta \sum_{j=1}^{k-1} j(k-j)a_j a_{k-j}$$

for $k \geq 2$. The generating function is formed by multiplying the k^{th} equation by kx^k and summing over all k . This gives

$$-x + \sum_{k=1}^{\infty} k a_k x^k + 2\delta x \sum_{k=1}^{\infty} a_k k^2 x^{k-1} = \delta \sum_{k=1}^{\infty} k x^k \sum_{j=1}^{k-1} j(k-j)a_j a_{k-j}.$$

Note that

$$g(x) = \sum_{k=1}^{\infty} k a_k x^k \text{ and } g'(x) = \sum_{k=1}^{\infty} a_k k^2 x^{k-1}.$$

Thus,

$$-x + g(x) + 2\delta x g'(x) = \delta \sum_{k=1}^{\infty} k x^k \sum_{j=1}^{k-1} j(k-j)a_j a_{k-j}.$$

Working with the right hand side

$$\delta \sum_{k=1}^{\infty} k x^k \sum_{j=1}^{k-1} j(k-j)a_j a_{k-j} = \delta x \sum_{k=1}^{\infty} \sum_{j=1}^{k-1} j(k-j)(j+k-j)x^{k-1} a_j a_{k-j}.$$

Now breaking the $j + k - j$ into two sums gives

$$\delta x \sum_{k=1}^{\infty} \sum_{j=1}^{k-1} j^2 a_j x^{j-1} (k-j) a_{k-j} x^{k-j} + \delta x \sum_{k=1}^{\infty} \sum_{j=1}^{k-1} j a_j x^j (k-j)^2 a_{k-j} x^{k-j-1}.$$

Notice that the second sum is obtained from the first by substituting $k-j$ for j and that both terms are $\delta x g'g$. Thus,

$$-x + g(x) + 2\delta x g'(x) = 2\delta x g'(x)g(x).$$

Hence,

$$g' = \frac{1}{2\delta} \frac{1 - \frac{g}{x}}{1 - g}.$$

Phase transition for nonfinite components

The generating function $g(x)$ contains information about the finite components of the graph. A finite component is a component of size $1, 2, \dots$ which does not depend on t . Observe that $g(1) = \sum_{k=0}^{\infty} ka_k$ and hence $g(1)$ is the probability that a randomly chosen vertex will belong to a component of finite size. If $g(1) = 1$ there are no nonfinite components. When $g(1) \neq 1$, then $1 - g(1)$ is the expected fraction of the vertices that are in nonfinite components. Potentially, there could be many such nonfinite components. But an argument similar to Part 3 of Theorem 3.9 concludes that two fairly large components would merge into one. Suppose there are two connected components at time t , each of size at least $t^{4/5}$. Consider the earliest created $\frac{1}{2}t^{4/5}$ vertices in each part. These vertices must have lived for at least $\frac{1}{2}t^{4/5}$ time after creation. At each time, the probability of an edge forming between two such vertices, one in each component, is at least $\delta\Omega(t^{-2/5})$ and so the probability that no such edge formed is at most $(1 - \delta t^{-2/5})^{t^{4/5}/2} \leq e^{-\Omega(\delta t^{2/5})} \rightarrow 0$. So with high probability, such components would have merged into one. But this still leaves open the possibility of many components of size t^ε , $(\ln t)^2$, or some other slowly growing function of t .

We now calculate the value of δ at which the phase transition for a nonfinite component occurs. Recall that the generating function for $g(x)$ satisfies

$$g'(x) = \frac{1}{2\delta} \frac{1 - \frac{g(x)}{x}}{1 - g(x)}.$$

If δ is greater than some $\delta_{critical}$, then $g(1) \neq 1$. In this case the above formula simplifies with $1 - g(1)$ canceling from the numerator and denominator, leaving just $\frac{1}{2\delta}$. Since ka_k is the probability that a randomly chosen vertex is in a component of size k , the average size of the finite components is $g'(1) = \sum_{k=1}^{\infty} k^2 a_k$. Now, $g'(1)$ is given by

$$g'(1) = \frac{1}{2\delta} \tag{3.6}$$

for all δ greater than $\delta_{critical}$. If δ is less than $\delta_{critical}$, then all vertices are in finite components. In this case $g(1) = 1$ and both the numerator and the denominator approach zero. Applying L'Hopital's rule

$$\lim_{x \rightarrow 1} g'(x) = \frac{1}{2\delta} \left. \frac{\frac{xg'(x) - g(x)}{x^2}}{g'(x)} \right|_{x=1}$$

or

$$(g'(1))^2 = \frac{1}{2\delta} (g'(1) - g(1)).$$

The quadratic $(g'(1))^2 - \frac{1}{2\delta}g'(1) + \frac{1}{2\delta}g(1) = 0$ has solutions

$$g'(1) = \frac{\frac{1}{2\delta} \pm \sqrt{\frac{1}{4\delta^2} - \frac{4}{2\delta}}}{2} = \frac{1 \pm \sqrt{1 - 8\delta}}{4\delta}. \quad (3.7)$$

The two solutions given by (3.7) become complex for $\delta > 1/8$ and thus can be valid only for $0 \leq \delta \leq 1/8$. For $\delta > 1/8$, the only solution is $g'(1) = \frac{1}{2\delta}$ and a nonfinite component exists. As δ is decreased, at $\delta = 1/8$ there is a singular point where for $\delta < 1/8$ there are three possible solutions, one from (3.6) which implies a giant component and two from (3.7) which imply no giant component. To determine which one of the three solutions is valid, consider the limit as $\delta \rightarrow 0$. In the limit all components are of size one since there are no edges. Only (3.7) with the minus sign gives the correct solution

$$g'(1) = \frac{1 - \sqrt{1 - 8\delta}}{4\delta} = \frac{1 - (1 - \frac{1}{2}8\delta - \frac{1}{4}64\delta^2 + \dots)}{4\delta} = 1 + 4\delta + \dots = 1.$$

In the absence of any nonanalytic behavior in the equation for $g'(x)$ in the region $0 \leq \delta < 1/8$, we conclude that (3.7) with the minus sign is the correct solution for $0 \leq \delta < 1/8$ and hence the critical value of δ for the phase transition is $1/8$. As we shall see, this is different from the static case.

As the value of δ is increased, the average size of the finite components increase from one to

$$\left. \frac{1 - \sqrt{1 - 8\delta}}{4\delta} \right|_{\delta=1/8} = 2$$

when δ reaches the critical value of $1/8$. At $\delta = 1/8$, the average size of the finite components jumps to $\frac{1}{2\delta}|_{\delta=1/8} = 4$ and then decreases as $\frac{1}{2\delta}$ as the giant component swallows up the finite components starting with the larger components.

Comparison to static random graph

Consider a static random graph with the same degree distribution as the graph in the growth model. Again let p_k be the probability of a vertex being of degree k . From (3.5)

$$p_k = \frac{(2\delta)^k}{(1 + 2\delta)^{k+1}}.$$

Recall the Molloy Reed analysis of random graphs with given degree distributions which asserts that there is a phase transition at $\sum_{i=0}^{\infty} i(i-2)p_i = 0$. Using this, it is easy to see that a phase transition occurs for $\delta = 1/4$. For $\delta = 1/4$,

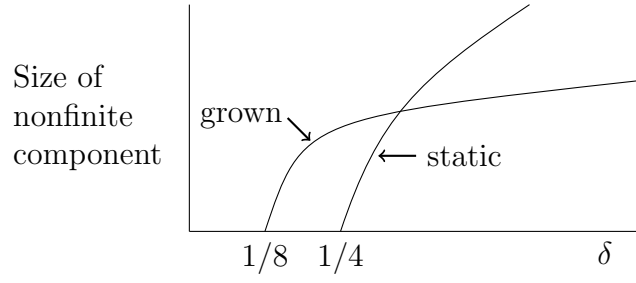


Figure 3.15: Comparison of the static random graph model and the growth model. The curve for the growth model is obtained by integrating g' .

$$p_k = \frac{(2\delta)^k}{(1+2\delta)^{k+1}} = \frac{\left(\frac{1}{2}\right)^k}{\left(1+\frac{1}{2}\right)^{k+1}} = \frac{\left(\frac{1}{2}\right)^k}{\frac{3}{2}\left(\frac{3}{2}\right)^k} = \frac{2}{3}\left(\frac{1}{3}\right)^k$$

and

$$\sum_{i=0}^{\infty} i(i-2)\frac{2}{3}\left(\frac{1}{3}\right)^i = \frac{2}{3}\sum_{i=0}^{\infty} i^2\left(\frac{1}{3}\right)^i - \frac{4}{3}\sum_{i=0}^{\infty} i\left(\frac{1}{3}\right)^i = \frac{2}{3} \times \frac{3}{2} - \frac{4}{3} \times \frac{3}{4} = 0.$$

Recall that $1 + a + a^2 + \dots = \frac{1}{1-a}$, $a + 2a^2 + 3a^3 + \dots = \frac{a}{(1-a)^2}$, and $a + 4a^2 + 9a^3 + \dots = \frac{a(1+a)}{(1-a)^3}$.

See references at end of the chapter for calculating the size S_{static} of the giant component in the static graph. The result is

$$S_{static} = \begin{cases} 0 & \delta \leq \frac{1}{4} \\ 1 - \frac{1}{\delta + \sqrt{\delta^2 + 2\delta}} & \delta > \frac{1}{4} \end{cases}$$

3.9.2 Growth Model With Preferential Attachment

Consider a growth model with preferential attachment. At each time unit, a vertex is added to the graph. Then with probability δ , an edge is attached to the new vertex and to a vertex selected at random with probability proportional to its degree. This model generates a tree with a power law distribution.

Let $d_i(t)$ be the expected degree of the i^{th} vertex at time t . The sum of the degrees of all vertices at time t is $2\delta t$ and thus the probability that an edge is connected to vertex i at time t is $\frac{d_i(t)}{2\delta t}$. The degree of vertex i is governed by the equation

$$\frac{\partial}{\partial t} d_i(t) = \delta \frac{d_i(t)}{2\delta t} = \frac{d_i(t)}{2t}$$

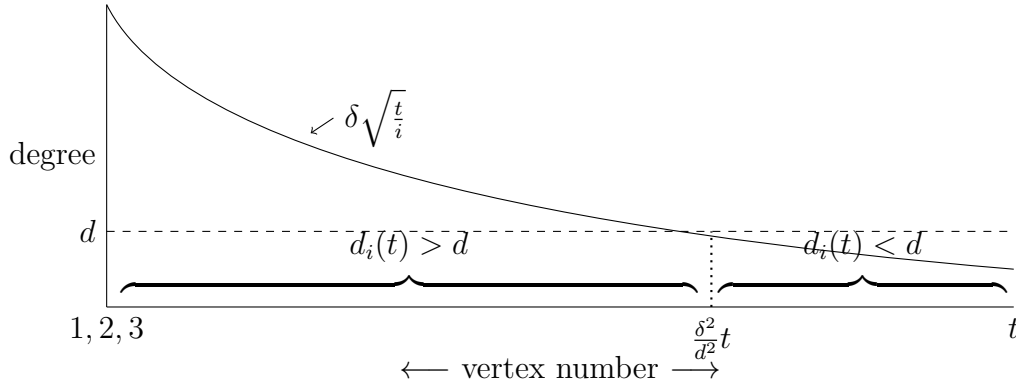


Figure 3.16: Illustration of degree of i^{th} vertex at time t . At time t , vertices numbered 1 to $\frac{\delta^2}{d^2}t$ have degrees greater than d .

where δ is the probability that an edge is added at time t and $\frac{d_i(t)}{2\delta t}$ is the probability that the vertex i is selected for the end point of the edge.

The two in the denominator governs the solution which is of the form $at^{\frac{1}{2}}$. The value of a is determined by the initial condition $d_i(t) = \delta$ at $t = i$. Thus, $\delta = ai^{\frac{1}{2}}$ or $a = \delta i^{-\frac{1}{2}}$. Hence, $d_i(t) = \delta \sqrt{\frac{t}{i}}$.

Next, we determine the probability distribution of vertex degrees. Now, $d_i(t)$ is less than d provided $i > \frac{\delta^2}{d^2}t$. The fraction of the t vertices at time t for which $i > \frac{\delta^2}{d^2}t$ and thus that the degree is less than d is $1 - \frac{\delta^2}{d^2}$. Hence, the probability that a vertex has degree less than d is $1 - \frac{\delta^2}{d^2}$. The probability density $P(d)$ satisfies

$$\int_0^d P(d) \partial d = \text{Prob}(\text{degree} < d) = 1 - \frac{\delta^2}{d^2}$$

and can be obtained from the derivative of $\text{Prob}(\text{degree} < d)$.

$$P(d) = \frac{\partial}{\partial d} \left(1 - \frac{\delta^2}{d^2} \right) = 2 \frac{\delta^2}{d^3},$$

a power law distribution.

3.10 Small World Graphs

In the 1960's, Stanley Milgram carried out an experiment that indicated that any two individuals in the United States were connected by a short sequence of acquaintances. Milgram would ask a source individual, say in Nebraska, to start a letter on its journey to a target individual in Massachusetts. The Nebraska individual would be given basic

information about the target including his address and occupation and asked to send the letter to someone he knew on a first name basis, who was closer to the target individual, in order to transmit the letter to the target in as few steps as possible. Each person receiving the letter would be given the same instructions. In successful experiments, it would take on average five to six steps for a letter to reach its target. This research generated the phrase “six degrees of separation” along with substantial research in social science on the interconnections between people. Surprisingly, there was no work on how to find the short paths using only local information.

In many situations, phenomena are modeled by graphs whose edges can be partitioned into local and long distance. We adopt a simple model of a directed graph due to Kleinberg, having local and long distance edges. Consider a 2-dimensional $n \times n$ grid where each vertex is connected to its four adjacent vertices. In addition to these local edges, there is one long distance edge out of each vertex. The probability that the long distance edge from vertex u terminates at v , $v \neq u$, is a function of the distance $d(u, v)$ from u to v . Here distance is measured by the shortest path consisting only of local grid edges. The probability is proportional to $1/d^r(u, v)$ for some constant r . This gives a one parameter family of random graphs. For r equal zero, $1/d^0(u, v) = 1$ for all u and v and thus the end of the long distance edge at u is uniformly distributed over all vertices independent of distance. As r increases the expected length of the long distance edge decreases. As r approaches infinity, there are no long distance edges and thus no paths shorter than that of the lattice path. What is interesting is that for r less than two, there are always short paths, but no local algorithm to find them. A local algorithm is an algorithm that is only allowed to remember the source, the destination, and its current location and can query the graph to find the long-distance edge at the current location. Based on this information, it decides the next vertex on the path.

The difficulty is that for $r < 2$, the end points of the long distance edges tend to be uniformly distributed over the vertices of the grid. Although short paths exist, it is unlikely on a short path to encounter a long distance edge whose end point is close to the destination. When r equals two, there are short paths and the simple algorithm that always selects the edge that ends closest to the destination will find a short path. For r greater than two, again there is no local algorithm to find a short path. Indeed, with high probability, there are no short paths at all.

The probability that the long distance edge from u goes to v is proportional to $d^{-r}(u, v)$. Note that the constant of proportionality will vary with the vertex u depending on where u is relative to the border of the $n \times n$ grid. However, the number of vertices at distance exactly k from u is at most $4k$ and for $k \leq n/2$ is at least k . Let $c_r(u) = \sum_v d^{-r}(u, v)$ be the normalizing constant. It is the inverse of the constant of proportionality.

$r > 2$ The lengths of long distance edges tend to be short so the probability of encountering a sufficiently long, long-distance edge is too low.

$r = 2$ Selecting the edge with end point closest to the destination finds a short path.

$r < 2$ The ends of long distance edges tend to be uniformly distributed. Short paths exist but a polylog length path is unlikely to encounter a long distance edge whose end point is close to the destination.

Figure 3.17: Effects of different values of r on the expected length of long distance edges and the ability to find short paths.

For $r > 2$, $c_r(u)$ is lower bounded by

$$c_r(u) = \sum_v d^{-r}(u, v) \geq \sum_{k=1}^{n/2} (k) k^{-r} = \sum_{k=1}^{n/2} k^{1-r} \geq 1.$$

No matter how large r is the first term of $\sum_{k=1}^{n/2} k^{1-r}$ is at least one.

For $r = 2$ the normalizing constant $c_r(u)$ is upper bounded by

$$c_r(u) = \sum_v d^{-r}(u, v) \leq \sum_{k=1}^{2n} (4k) k^{-2} \leq 4 \sum_{k=1}^{2n} \frac{1}{k} = \theta(\ln n).$$

For $r < 2$, the normalizing constant $c_r(u)$ is lower bounded by

$$c_r(u) = \sum_v d^{-r}(u, v) \geq \sum_{k=1}^{n/2} (k) k^{-r} \geq \sum_{k=n/4}^{n/2} k^{1-r}.$$

The summation $\sum_{k=n/4}^{n/2} k^{1-r}$ has $\frac{n}{4}$ terms, the smallest of which is $(\frac{n}{4})^{1-r}$ or $(\frac{n}{2})^{1-r}$ depending on whether r is greater or less than one. This gives the following lower bound on $c_r(u)$.

$$c_r(u) \geq \frac{n}{4} \omega(n^{1-r}) = \omega(n^{2-r}).$$

No short paths exist for the $r > 2$ case.

For $r > 2$, we first show that for at least one half the pairs of vertices there is no short path between them. We begin by showing that the expected number of edges of length

greater than $n^{\frac{r+2}{2r}}$ goes to zero. The probability of an edge from u to v is $d^{-r}(u, v)/c_r(u)$ where $c_r(u)$ is lower bounded by a constant. Thus, the probability that a long edge is of length greater than or equal to $n^{\frac{r+2}{2r}}$ is upper bounded by some constant c times $\left(n^{\frac{r+2}{2r}}\right)^{-r}$ or $cn^{-(\frac{r+2}{2})}$. Since there are n^2 long edges, the expected number of edges of length at least $n^{\frac{r+2}{2r}}$ is at most $cn^2n^{-(\frac{r+2}{2})}$ or $cn^{\frac{2-r}{2}}$, which for $r > 2$ goes to zero. Thus, by the first moment method, almost surely, there are no such edges.

For at least one half of the pairs of vertices, the grid distance, measured by grid edges between the vertices, is greater than or equal to $n/4$. Any path between them must have at least $\frac{1}{4}n/n^{\frac{r+2}{2r}} = \frac{1}{4}n^{\frac{r-2}{2r}}$ edges since there are no edges longer than $n^{\frac{r+2}{2r}}$ and so there is no polylog length path.

An algorithm for the $r = 2$ case

For $r = 2$, the local algorithm that selects the edge that ends closest to the destination t finds a path of expected length $O(\ln n)^3$. Suppose the algorithm is at a vertex u which is at distance k from t . Then within an expected $O(\ln n)^2$ steps, the algorithm reaches a point at distance at most $k/2$. The reason is that there are $\Omega(k^2)$ vertices at distance at most $k/2$ from t . Each of these vertices is at distance at most $k + k/2 = O(k)$ from u . See Figure 3.18. Recall that the normalizing constant c_r is upper bounded by $O(\ln n)$, and hence, the constant of proportionality is lower bounded by some constant times $1/\ln n$. Thus, the probability that the long-distance edge from u goes to one of these vertices is at least

$$\Omega(k^2k^{-r}/\ln n) = \Omega(1/\ln n).$$

Consider $\Omega(\ln n)^2$ steps of the path from u . The long-distance edges from the points visited at these steps are chosen independently and each has probability $\Omega(1/\ln n)$ of reaching within $k/2$ of t . The probability that none of them does is

$$(1 - \Omega(1/\ln n))^{c(\ln n)^2} = c_1 e^{-\ln n} = \frac{c_1}{n}$$

for a suitable choice of constants. Thus, the distance to t is halved every $O(\ln n)^2$ steps and the algorithm reaches t in an expected $O(\ln n)^3$ steps.

A local algorithm cannot find short paths for the $r < 2$ case

For $r < 2$ no local polylog time algorithm exists for finding a short path. To illustrate the proof, we first give the proof for the special case $r = 0$, and then give the proof for $r < 2$.

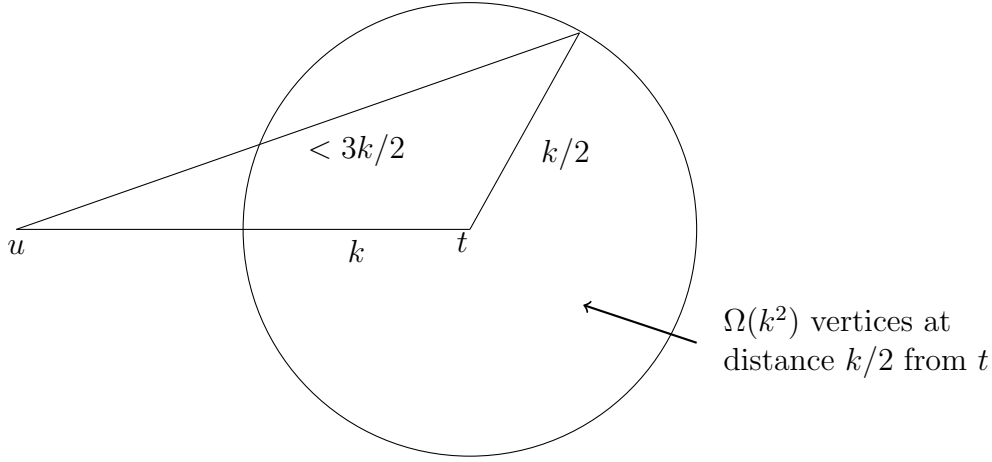


Figure 3.18: Small worlds.

When $r = 0$, all vertices are equally likely to be the end point of a long distance edge. Thus, the probability of a long distance edge hitting one of the n vertices that are within distance $\sqrt{n}$ of the destination is $1/n$. Along a path of length $\sqrt{n}$, the probability that the path does not encounter such an edge is $(1 - 1/n)^{\sqrt{n}}$. Now,

$$\lim_{n \rightarrow \infty} \left(1 - \frac{1}{n}\right)^{\sqrt{n}} = \lim_{n \rightarrow \infty} \left(1 - \frac{1}{n}\right)^{n \frac{1}{\sqrt{n}}} = \lim_{n \rightarrow \infty} e^{-\frac{1}{\sqrt{n}}} = 1.$$

Since with probability $1/2$ the starting point is at distance at least $n/4$ from the destination and in $\sqrt{n}$ steps, the path will not encounter a long distance edge ending within distance $\sqrt{n}$ of the destination, for at least half of the starting points the path length will be at least $\sqrt{n}$. Thus, the expected time is at least $\frac{1}{2}\sqrt{n}$ and hence not in polylog time.

For the general $r < 2$ case, we show that a local algorithm cannot find paths of length $O(n^{(2-r)/4})$. Let $\delta = (2 - r)/4$ and suppose the algorithm finds a path with at most n^δ edges. There must be a long-distance edge on the path which terminates within distance n^δ of t ; otherwise, the path would end in n^δ grid edges and would be too long. There are $O(n^{2\delta})$ vertices within distance n^δ of t and the probability that the long distance edge from one vertex of the path ends at one of these vertices is at most $n^{2\delta} \left(\frac{1}{n^{2-r}}\right) = n^{(r-2)/2}$. To see this, recall that the lower bound on the normalizing constant is $\theta(n^{2-r})$ and hence an upper bound on the probability of a long distance edge hitting v is $\theta\left(\frac{1}{n^{2-r}}\right)$ independent of where v is. Thus, the probability that the long distance edge from one of the n^δ vertices on the path hits any one of the $n^{2\delta}$ vertices within distance n^δ of t is $n^{2\delta} \frac{1}{n^{2-r}} = n^{\frac{r-2}{2}}$. The probability that this happens for any one of the n^δ vertices on the path is at most $n^{\frac{r-2}{2}} n^\delta = n^{\frac{r-2}{2}} n^{\frac{2-r}{4}} = n^{(r-2)/4} = o(1)$ as claimed.

Short paths exist for $r < 2$

Finally we show for $r < 2$ that there are $O(\ln n)$ length paths between s and t . The proof is similar to the proof of Theorem 3.17 showing $O(\ln n)$ diameter for $G(n, p)$ when p is $\Omega(\ln n/n)$, so we do not give all the details here. We give the proof only for the case when $r = 0$.

For a particular vertex v , let S_i denote the set of vertices at distance i from v . Using only local edges, if i is $O(\sqrt{\ln n})$, then $|S_i|$ is $\Omega(\ln n)$. For later i , we argue a constant factor growth in the size of S_i as in Theorem 3.17. As long as $|S_1| + |S_2| + \dots + |S_i| \leq n^2/2$, for each of the $n^2/2$ or more vertices outside, the probability that the vertex is not in S_{i+1} is $(1 - \frac{1}{n^2})^{|S_i|} \leq 1 - \frac{|S_i|}{2n^2}$ since the long-distance edge from each vertex of S_i chooses a long-distance neighbor at random. So, the expected size of S_{i+1} is at least $|S_i|/4$ and using Chernoff, we get constant factor growth up to $n^2/2$. Thus, for any two vertices v and w , the number of vertices at distance $O(\ln n)$ from each is at least $n^2/2$. Any two sets of cardinality at least $n^2/2$ must intersect giving us a $O(\ln n)$ length path from v to w .

3.11 Bibliographic Notes

The $G(n, p)$ random graph model is from Erdős Rényi [ER60]. Among the books written on properties of random graphs a reader may wish to consult Palmer [Pal85], Jansen, Luczak and Ruciński [JLR00], or Bollobás [Bol01]. Material on phase transitions can be found in [BT87]. The work on phase transitions for CNF was started by Chao and Franco [CF86]. Further work was done in [FS96], [AP03], [Fri99], and others. The proof here that the SC algorithm produces a solution when the number of clauses is cn for $c < \frac{2}{3}$ is from [Chv92].

For material on the giant component consult [Kar90] or [JKLP93]. Material on branching process can be found in [AN72]. The phase transition for giant components in random graphs with given degree distributions is from Molloy and Reed [MR95a].

There are numerous papers on growth models. The material in this chapter was based primarily on [CHK⁺] and [BA]. The material on small world is based on Kleinberg, [Kle00] which follows earlier work by Watts and Strogatz [WS98a].

3.12 Exercises

Exercise 3.1 Search the World Wide Web to find some real world graphs in machine readable form or data bases that could automatically be converted to graphs.

1. Plot the degree distribution of each graph.
2. Compute the average degree of each graph.
3. Count the number of connected components of each size in each graph.
4. Describe what you find.

Exercise 3.2 Find a data base in machine readable form that can be viewed as a graph. What is the expected number of edges? If the graph were a $G(n, p)$ graph, what would the value of p be? Find the number of components of various sizes. Check that your work is correct by multiplying the number of components of size s by s and summing over all sizes. Is the sum equal to the total number of vertices? Examine the small components and see if any have cycles.

Exercise 3.3 In $G(n, p)$ the probability of a vertex having degree k is $\binom{n}{k} p^k (1 - p)^{n-k}$.

1. Show by direct calculation that the expected degree is np .
2. Compute directly the variance of the distribution.
3. Where is the mode of the binomial distribution for a given value of p ? The mode is the point at which the probability is maximum.

Exercise 3.4

1. Plot the degree distribution for $G(1000, 0.003)$.
2. Plot the degree distribution for $G(1000, 0.030)$.

Exercise 3.5 In $G(n, \frac{1}{n})$, what is the probability that there is a vertex of degree $\log n$? Give an exact formula; also derive simple approximations.

Exercise 3.6 The example of Section 3.1.1 showed that if the degrees in $G(n, \frac{1}{n})$ were independent there would almost surely be a vertex of degree $\log n / \log \log n$. However, the degrees are not independent. Show how to overcome this difficulty.

Exercise 3.7 Let $f(n)$ be a function that is asymptotically less than n . Some such functions are $1/n$, a constant d , $\log n$ or $n^{\frac{1}{3}}$. Show that

$$\left(1 + \frac{f(n)}{n}\right)^n \simeq e^{f(n)}.$$

for large n . That is

$$\lim_{n \rightarrow \infty} \frac{\left(1 + \frac{f(n)}{n}\right)^n}{e^{f(n)}} = 1.$$

Exercise 3.8 Simplify $(1 - \frac{1}{n})^{n \log n}$.

Exercise 3.9 Generate a graph $G(n, \frac{d}{n})$ with $n = 1000$ and $d=2, 3$, and 6 . Count the number of triangles in each graph. Try the experiment with $n=100$.

Exercise 3.10 What is the expected number of squares (4-cycles) in $G(n, \frac{d}{n})$? What is the expected number of 4-cliques in $G(n, \frac{d}{n})$?

Exercise 3.11 Carry out an argument, similar to the one used for triangles, to show that $p = \frac{1}{n^{2/3}}$ is a threshold for the existence of a 4-clique. A 4-clique consists of four vertices with all $\binom{4}{2}$ edges present.

Exercise 3.12 What is the expected number of paths of length 3, $\log n$, $\sqrt{n}$, and $n - 1$ in $G(n, \frac{d}{n})$? The expected number of paths of a given length being infinite does not imply that a graph selected at random has such a path.

Exercise 3.13 Consider $G(n, \frac{1}{2})$. Give an algorithm that with high probability will find

1. a clique of size $\log n$.
2. an independent set of size $\log n$. A set of vertices is an independent set if there is no edge between any pair of vertices in the set.
3. a subgraph S in $G(n, \frac{1}{2})$, where S is any specified graph with $\log n$ vertices.

Exercise 3.14 Let x be an integer chosen uniformly at random from $\{1, 2, \dots, n\}$. Count the number of distinct prime factors of n . The exercise is to show that the number of prime factors almost surely is $\Theta(\ln \ln n)$. Let p stand for a prime number between 2 and n .

1. For each fixed prime p , let I_p be the indicator function of the event that p divides x . Show that $E(I_p) = \frac{1}{p} + O(\frac{1}{n})$. It is known that $\sum_{p \leq n} \frac{1}{p} = \ln \ln n$ and you may assume this.
2. The random variable of interest, $y = \sum_p I_p$, is the number of prime divisors of x picked at random. Show that the variance of y is $O(\ln \ln n)$. For this, assume the known result that the number of primes up to n is $O(n / \ln n)$. To bound the variance of y , think of what $E(I_p I_q)$ is for $p \neq q$, both primes.
3. Use (1) and (2) to prove that the number of prime factors is almost surely $\theta(\ln \ln n)$.

Exercise 3.15 Show for $\epsilon > 0$ that with high probability there exists a clique of size $(2 - \epsilon) \log n$ in $G(n, \frac{1}{2})$, but no clique of size $2 \log n$.

Exercise 3.16 Suppose one hides a clique of size k in a random graph $G(n, \frac{1}{2})$. I.e., in the random graph, choose some subset S of k vertices and put in the missing edges to make S a clique. Presented with the modified graph, find S . The larger S is, the easier it should be to find. In fact, if k is more than $c\sqrt{n \ln n}$, then the clique leaves a telltale sign identifying S as the k vertices of largest degree. Prove this statement by appealing to Theorem 3.1.1. It remains a puzzling open problem to do this when k is smaller, say, $O(n^{1/3})$.

Exercise 3.17 The clique problem in a graph is to find the maximal size clique. This problem is known to be NP-hard and so a polynomial time algorithm is thought unlikely. We can ask the corresponding question about random graphs. For example, in $G(n, \frac{1}{2})$ there almost surely is a clique of size $(2 - \epsilon) \log n$ for any $\epsilon > 0$. But it is not known how to find one in polynomial time.

1. Show that in $G(n, \frac{1}{2})$, there are, almost surely, no cliques of size $2 \log_2 n$.
2. Use the second moment method to show that in $G(n, \frac{1}{2})$, almost surely there are cliques of size $(2 - \epsilon) \log_2 n$.
3. Show that for any $\epsilon > 0$, a clique of size $(2 - \epsilon) \log n$ can be found in $G(n, \frac{1}{2})$ in time $n^{O(\ln n)}$.
4. Give an $O(n^2)$ algorithm for finding a clique of size $\Omega(\log n)$ in $G(n, \frac{1}{2})$. Hint: use a greedy algorithm. Apply your algorithm to $G(1000, \frac{1}{2})$. What size clique do you find?
5. An independent set of vertices in a graph is a set of vertices, no two of which are connected by an edge. Give a polynomial time algorithm for finding an independent set in $G(n, \frac{1}{2})$ of size $\Omega(\log n)$.

Exercise 3.18 Does there exist a copy of every subgraph with $(2 - \epsilon) \log n$ vertices and $\frac{1}{4} \binom{(2 - \epsilon) \log n}{2}$ edges in $G(n, \frac{1}{4})$?

Exercise 3.19 Given two instances, G_1 and G_2 of $G(n, \frac{1}{2})$, what is the largest subgraph common to both G_1 and G_2 ?

Exercise 3.20 (Birthday problem) What is the number of integers that must be drawn with replacement from a set of n integers so that some integer, almost surely, will be selected twice?

Exercise 3.21 Suppose the graph of a social network has 20,000 vertices. You have a program that starting from a random seed produces a community. A community is a set of vertices where each vertex in the set has more edges connecting it to other vertices in the set than to vertices outside of the set. In running the algorithm you find thousands of communities and wonder how many communities there are in the graph. Finally, when you find the 10,000th community, it is a duplicate. It is the same community as one found earlier.

1. Use the birthday problem to derive a lower bound on the number of communities.
2. Why do you only get a lower bound and not a good estimate?

Exercise 3.22 To better understand the binomial distribution plot $\binom{n}{k}p^k(1-p)^{n-k}$ as a function of k for $n = 50$ and $k = 0.05, 0.5, 0.95$. For each value of p check the sum over all k to ensure that the sum is one.

Exercise 3.23 Consider the binomial distribution $\binom{n}{i} \left(1 - \left(1 - \frac{d}{n}\right)^i\right)$ for $d > 1$. Here the distribution giving the probability of drawing i items is a different distribution for each value of i . Prove that as $n \rightarrow \infty$, the distribution goes to zero for all i except for i in the two ranges $[0, c_1 \log n]$ and $[\theta n - c_2 \sqrt{n}, \theta n + c_2 \sqrt{n}]$.

Exercise 3.24 Let s be the expected number of vertices discovered as a function of the number of steps t in a breadth first search of $G(n, \frac{d}{n})$. Write a differential equation using expected values for the size of s . Show that the normalized size $f = \frac{s-t}{n}$ of the frontier is $f(x) = 1 - e^{-dx} - x$ where $x = \frac{t}{n}$ is the normalized time.

Exercise 3.25 The normalized frontier in a breadth first search of $G(n, \frac{d}{n})$ is $f(x) = 1 - e^{-dx} - x$. For $d > 1$ let θ be the unique root in $(0, 1)$ of $1 - e^{-dx} - x = 0$. Prove that the expected value of the size of the frontier increases varies with i for i in the neighborhood of θ .

Exercise 3.26 For $f(x) = 1 - e^{-dx} - x$, what is the value of $x_{\max} = \arg \max f(x)$? What is the value of $f(x_{\max})$? Where does the maximum expected value of the frontier of a breadth search in $G(n, \frac{d}{n})$ occur as a function of n ?

Exercise 3.27 If y and z are independent, nonnegative random variables, then the generating function of the sum $y + z$ is the product of the generating function of y and z . Show that this follows from $E(x^{y+z}) = E(x^y x^z) = E(x^y)E(x^z)$.

Exercise 3.28 Let $f_j(x)$ be the j^{th} iterate of the generating function $f(x)$ of a branching process. When $m > 1$, $\lim_{j \rightarrow \infty} f_j(x) = q$ for $0 < x < 1$. In the limit this implies $\text{Prob}(z_j = 0) = q$ and $\text{Prob}(z_j = i) = 0$ for all nonzero finite values of i . Shouldn't the probabilities add up to 1? Why is this not a contradiction?

Exercise 3.29 Try to create a probability distribution for a branching process which varies with the current population in which future generations neither die out, nor grow to infinity.

Exercise 3.30 Let d be a constant strictly greater than 1. Show that for a branching process with number of children distributed as $\text{Binomial}(n - c_1 n^{2/3}, \frac{d}{n})$, the root of the $f(x) = 1$ in $(0, 1)$ is at most a constant strictly less than 1.

Exercise 3.31 Randomly generate $G(50, p)$ for several values of p . Start with $p = \frac{1}{50}$.

1. For what value of p do cycles first appear?
2. For what value of p do isolated vertices disappear and the graphs become connected?

Exercise 3.32 Consider $G(n, p)$ with $p = \frac{1}{3n}$. Then, almost surely, there are no cycles of length 10.

1. Use the second moment method to show that, almost surely, there is a simple path of length 10.
2. What goes wrong if we try to modify the argument that, almost surely, there are no cycles of length 10 to show that there is no path of length 10?

Exercise 3.33 Complete the second moment argument of Theorem 3.13 to show that for $p = \frac{d}{n}$, $d > 1$, $G(n, p)$ almost surely has a cycle.

Hint: If two cycles share one or more edges, then the union of the two cycles is at least one greater than the union of the vertices.

Exercise 3.34 Let $G(n, p)$ be a random graph and let x be the random variable denoting the number of unordered pairs of nonadjacent vertices (u, v) such that no other vertex of G is adjacent to both u and v . Prove that if $\lim_{n \rightarrow \infty} E(x) = 0$, then for large n there are almost no disconnected graphs, i.e. $\text{Prob}(x = 0) \rightarrow 1$ and hence $\text{Prob}(G \text{ is connected}) \rightarrow 1$. Actually, the graph becomes connected long before this condition is true.

Exercise 3.35 Draw a tree with 10 vertices and label each vertex with a unique integer from 1 to 10. Construct the Prüfer sequence (Appendix 11.8.7) for the tree. Given the Prüfer sequence, recreate the tree.

Exercise 3.36 Construct the tree corresponding to the following Prüfer sequences (Appendix 11.8.7)

1. 113663 (1,2), (1,3), (1,4), (3,5), (3,6), (6,7), and (6,8)
2. 552833226.

Exercise 3.37 What is the expected number of isolated vertices in $G(n, p)$ for $p = \frac{1}{2} \frac{\ln n}{n}$?

Exercise 3.38 Theorem 3.17 shows that for some $c > 0$ and $p = c \ln n / n$, $G(n, p)$ has diameter $O(\ln n)$. Tighten the argument to pin down as low a value as possible for c .

Exercise 3.39 Let $f(n)$ be a function that is asymptotically less than n . Some such functions are $1/n$, a constant d , $\log n$ or $n^{\frac{1}{3}}$. Show that

$$\left(1 + \frac{f(n)}{n}\right)^n \simeq e^{f(n)}.$$

for large n . That is

$$\lim_{n \rightarrow \infty} \frac{\left(1 + \frac{f(n)}{n}\right)^n}{e^{f(n)}} = 1.$$

Exercise 3.40 What is diameter of $G(n, p)$ for various values of p ?

Exercise 3.41

1. List five increasing properties of $G(n, p)$.
2. List five non increasing properties .

Exercise 3.42 Consider generating the edges of a random graph by flipping two coins, one with probability p_1 of heads and the other with probability p_2 of heads. Add the edge to the graph if either coin comes down heads. What is the value of p for the generated $G(n, p)$ graph?

Exercise 3.43 In the proof of Theorem 3.19, we proved for $p_0(n)$ such that $\lim_{n \rightarrow \infty} \frac{p_0(n)}{p(n)} = 0$ that $G(n, p_0)$ almost surely did not have property Q . Give the symmetric argument that for any $p_1(n)$ such that $\lim_{n \rightarrow \infty} \frac{p_1(n)}{p(n)} = 0$, $G(n, p_1)$ almost surely has property Q .

Exercise 3.44 Consider a model of a random subset $N(n, p)$ of integers $\{1, 2, \dots, n\}$ where, $N(n, p)$ is the set obtained by independently at random including each of $\{1, 2, \dots, n\}$ into the set with probability p . Define what an “increasing property” of $N(n, p)$ means. Prove that every increasing property of $N(n, p)$ has a threshold.

Exercise 3.45 $N(n, p)$ is a model of a random subset of integers $\{1, 2, \dots, n\}$ where, $N(n, p)$ is the set obtained by independently at random including each of $\{1, 2, \dots, n\}$ into the set with probability p . What is the threshold for $N(n, p)$ to contain

1. a perfect square,
2. a perfect cube,
3. an even number,
4. three numbers such that $x + y = z$?

Exercise 3.46 Explain why the property, that $N(n, p)$ contains the integer 1, has a threshold. What is the threshold?

Exercise 3.47 Is there a condition such that any property satisfying the condition has a sharp threshold? For example, is monotonicity such a condition?

Exercise 3.48 The Sudoku game consists of a 9×9 array of squares. The array is partitioned into nine 3×3 squares. Each small square should be filled with an integer between 1 and 9 so that each row, each column, and each 3×3 square contains exactly one copy of each integer. Initially the board has some of the small squares filled in in such a way that there is exactly one way to complete the assignments of integers to squares. Some simple rules can be developed to fill in the remaining squares such as if the row and column

containing a square already contain a copy of every integer except one, that integer should be placed in the square.

Start with a 9×9 array of squares with each square containing a number between 1 and 9 such that no row, column, or 3×3 square has two copies of any integer.

1. How many integers can you randomly erase and there still be only one way to correctly fill in the board?
2. Develop a set of simple rules for filling in squares such as if a row does not contain a given integer and if every column except one in which the square in the row is blank contains the integer, then place the integer in the remaining blank entry in the row. How many integers can you randomly erase and your rules will still completely fill in the board?

Exercise 3.49 Generalize the Sudoku game for arrays of size $n^2 \times n^2$. Develop a simple set of rules for completing the game. An example of a rule is the following. If the a row does not contain a given integer and if every column except one in which the square in the row is blank contains the integer, then place the integer in the remaining blank entry in the row. Start with a legitimate completed array and erase k entries at random.

1. Is there a threshold for the integer k such that if only k entries of the array are erased, your set of rules will find a solution?
2. Experimentally determine k for some large value of n .

Exercise 3.50 Let $\{x_i | 1 \leq i \leq n\}$, be a set of indicator variables with identical probability distributions. Let $x = \sum_{i=1}^n x_i$ and suppose $E(x) \rightarrow \infty$. Show that if the x_i are statistically independent, then $\text{Prob}(x = 0) \rightarrow 0$.

Exercise 3.51 In a square $n \times n$ grid, each of the $O(n^2)$ edges is randomly chosen to be present with probability p and absent with probability $1 - p$. Consider the increasing property that there is a path from the bottom left corner to the top right corner which always goes to the right or up. Show that $p = 1/2$ is a threshold for the property. Is it a sharp threshold?

Exercise 3.52 The threshold property seems to be related to uniform distributions. What if we considered other distributions? Consider a model where i is selected from the set $\{1, 2, \dots, n\}$ with probability $\frac{c(n)}{i}$. Is there a threshold for perfect squares? Is there a threshold for arithmetic progressions?

Exercise 3.53 Modify the proof that every increasing property of $G(n, p)$ has a threshold to apply to the 3-CNF satisfiability problem.

Exercise 3.54 Evaluate $(1 - \frac{1}{2^k})^{2^k}$ for $k=3, 5$, and 7 . How close is it to $1/e$?

Exercise 3.55 Randomly generate clauses for a Boolean formula in 3-CNF. Compute the number of solutions and the number of connected components of the solution set as a function of the number of clauses generated. What happens?

Exercise 3.56 Consider a random process for generating a Boolean function f in conjunctive normal form where each of c clauses is generated by placing each of n variables in the clause with probability p and complementing the variable with probability $1/2$. What is the distribution of clause sizes for various p such as $p = 3/n$, $1/2$, other values? Experimentally determine the threshold value of p for f to cease to be satisfied.

Exercise 3.57 For a random 3-CNF formula with n variables and cn clauses, what is the expected number of satisfying assignments?

Exercise 3.58 Which of the following variants of the SC algorithm admit a theorem like Theorem 3.21?

1. Among all clauses of least length, pick the first one in the order in which they appear in the formula.
2. Set the literal appearing in most clauses independent of length to 1.

Exercise 3.59 Suppose we have a queue of jobs serviced by one server. There is a total of n jobs in the system. At time t , each remaining job independently decides to join the queue to be serviced with probability $p = d/n$, where $d < 1$ is a constant. Each job has a processing time of 1 and at each time the server services one job, if the queue is nonempty. Show that with high probability, no job waits more than $\Omega(\ln n)$ time to be serviced once it joins the queue.

Exercise 3.60 Consider $G(n, p)$.

1. Where is phase transition for 2-colorability? Hint: For $p = d/n$ with $d < 1$, $G(n, p)$ is acyclic, so it is bipartite and hence 2-colorable. When $pn \rightarrow \infty$, the expected number of triangles goes to infinity. Show that, almost surely, there is a triangle? What does this do for 2-colorability?
2. What about 3-colorability?

Exercise 3.61 A vertex cover of size k for a graph is a set of k vertices such that one end of each edge is in the set. Experimentally play with the following problem. For $G(n, \frac{1}{2})$, for what value of k is there a vertex cover of size k ?

Exercise 3.62 Consider graph 3-colorability. Randomly generate the edges of a graph and compute the number of solutions and the number of connected components of the solution set as a function of the number of edges generated. What happens?

Exercise 3.63 In $G(n, p)$, let x_k be the number of connected components of size k . Using x_k , write down the probability that a randomly chosen vertex is in a connected component of size k . Also write down the expected size of the connected component containing a randomly chosen vertex.

Exercise 3.64 For p asymptotically greater than $\frac{1}{n}$, show that

$$\sum_{i=0}^{\infty} i(i-2)\lambda_i > 0.$$

Exercise 3.65 Consider generating a random graph adding one edge at a time. Let $n(i, t)$ be the number of components of size i at time t .

$$\begin{aligned} n(1, 1) &= n \\ n(1, t) &= 0 \quad t > 1 \\ n(i, t) &= n(i, t-1) + \sum \frac{j(i-j)}{n^2} n(j, t-1) n(i-j, t-1) - \frac{2i}{n} n(i) \end{aligned}$$

Compute $n(i, t)$ for a number of values of i and t . What is the behavior? What is the sum of $n(i, t)$ for fixed t and all i ? Can you write a generating function for $n(i, t)$?

Exercise 3.66 The global clustering coefficient of a graph is defined as follows. Let d_v be the degree of vertex v and let e_v be the number of edges connecting vertices adjacent to vertex v . The global clustering coefficient c is given by

$$c = \sum_v \frac{2e_v}{d_v(d_v-1)}.$$

In a social network, for example, it measures what fraction of pairs of friends of each person are themselves friends. If many are, the clustering coefficient is high. What is c for a random graph with $p = \frac{d}{n}$? For a denser graph? Compare this value to that for some social network.

Exercise 3.67 Consider a structured graph, such as a grid or cycle, and gradually add edges or reroute edges at random. Let L be the average distance between all pairs of vertices in a graph and let C be the ratio of triangles to connected sets of three vertices. Plot L and C as a function of the randomness introduced.

Exercise 3.68 Consider an $n \times n$ grid in the plane.

1. Prove that for any vertex u , there are at least k vertices at distance k for $1 \leq k \leq n/2$.
2. Prove that for any vertex u , there are at most $4k$ vertices at distance k .
3. Prove that for one half of the pairs of points, the distance between them is at least $4/4$.

Exercise 3.69 Show that in a small-world graph with $r \leq 2$, that there exist short paths with high probability. The proof for $r = 0$ is in the text.

Exercise 3.70 Change the small worlds graph as follows. Start with a $n \times n$ grid where each vertex has one long-distance edge to a vertex chosen uniformly at random. These are exactly like the long-distance edges for $r = 0$. But the grid edges are not present. Instead, we have some other graph with the property that for each vertex, there are $\Theta(t^2)$ vertices at distance t from the vertex for $t \leq n$. Show that, almost surely, the diameter is $O(\ln n)$.

Exercise 3.71 Given an n node directed graph with two random out edges from each node. For two vertices s and t chosen at random, prove that there exists a path of length at most $O(\ln n)$ from s to t with high probability.

Exercise 3.72 How does the diameter of a graph consisting of a cycle change as one adds a few random long distance edges? This question explores how much randomness is needed to get a small world.

Exercise 3.73 Ideas and diseases spread rapidly in small world graphs. What about spread of social contagion? A disease needs only one contact and with some probability transfers. Social contagion needs several contacts. How many vertices must one start with to spread social contagion, if the spread of contagion requires two adjacent vertices?

Exercise 3.74 How many edges are needed to disconnect a small world graph? By disconnect we mean at least two pieces each of reasonable size. Is this connected to the emergence of a giant component?

Exercise 3.75 In the small world model, would it help if the algorithm could look at edges at any node at a cost of one for each node looked at?

Exercise 3.76 Consider the $n \times n$ grid in the section on small world graphs. If the probability of an edge from vertex u to vertex v is proportional to $d^{-r}(u, v)$, show that the constant of proportionality $c_r(u)$ is

$$\begin{array}{ll} \theta(n^{2-r}) & \text{for } r > 2 \\ \theta(\ln n) & \text{for } r = 2 \\ \theta(1) & \text{for } r < 2 \end{array}$$

Exercise 3.77 In the $n \times n$ grid prove that for at least half of the pairs of vertices, the distance between the vertices is greater than or equal to $n/4$

Exercise 3.78 Show that for $r < 2$ in the small world graph model that short paths exist but a polylog length path is unlikely to encounter a long distance edge whose end point is close to the destination.

Exercise 3.79 Make a list of the ten most interesting things you learned about random graphs.

4 Singular Value Decomposition (SVD)

The singular value decomposition of a matrix A is the factorization of A into the product of three matrices, $A = UDV^T$, where the columns of U and V are orthonormal and the matrix D is diagonal with positive real entries. The SVD is useful in many tasks. In many applications, a data matrix A is close to a matrix of low rank and it is useful to find a good low rank approximation to A . The singular value decomposition of A gives one the best rank k approximation to A , for any k .

The singular value decomposition is defined for all matrices, whereas the more commonly used eigenvector decomposition requires the matrix A be square and certain other conditions on the matrix to ensure orthogonality of the eigenvectors. In contrast, the columns of V in the singular value decomposition, called the *right-singular vectors* of A , always form an orthogonal set with no assumptions on A . The columns of U are called the *left-singular vectors* and they also form an orthogonal set. A simple consequence of the orthonormality is that for a square and invertible matrix A , the inverse of A is $VD^{-1}U^T$.

To gain insight into the SVD, treat the rows of an $n \times d$ matrix A as n points in a d -dimensional space and consider the problem of finding the best k -dimensional subspace with respect to the set of points. Here best means minimize the sum of the squares of the perpendicular distances of the points to the subspace. We begin with a special case where the subspace is 1-dimensional, namely a line through the origin. We will see later that the best-fitting k -dimensional subspace can be found by k applications of the best fitting line algorithm. The problem is called the *best least squares fit*.

Consider projecting a point $\mathbf{a}_i = (a_{i1}, a_{i2}, \dots, a_{id})$ onto a line through the origin. Then

$$a_{i1}^2 + a_{i2}^2 + \dots + a_{id}^2 = (\text{length of projection})^2 + (\text{distance of point to line})^2.$$

See Figure 4.1. Thus

$$(\text{distance of point to line})^2 = a_{i1}^2 + a_{i2}^2 + \dots + a_{id}^2 - (\text{length of projection})^2.$$

Since $\sum_{i=1}^n (a_{i1}^2 + a_{i2}^2 + \dots + a_{id}^2)$ is a constant independent of the line, minimizing the sum of the squares of the distances to the line is equivalent to maximizing the sum of the squares of the lengths of the projections onto the line. Similarly for best-fit subspaces, maximizing the sum of the squared lengths of the projections onto the subspace minimizes the sum of squared distances to the subspace.

The reader may wonder why we minimize the sum of squared perpendicular distances to the line. We could alternatively have defined the best-fit line to be the one that minimizes the sum of perpendicular distances to the line. There are examples where this definition gives a different answer than the line minimizing the sum of squared perpendicular distances. The choice of the objective function as the sum of squared distances seems

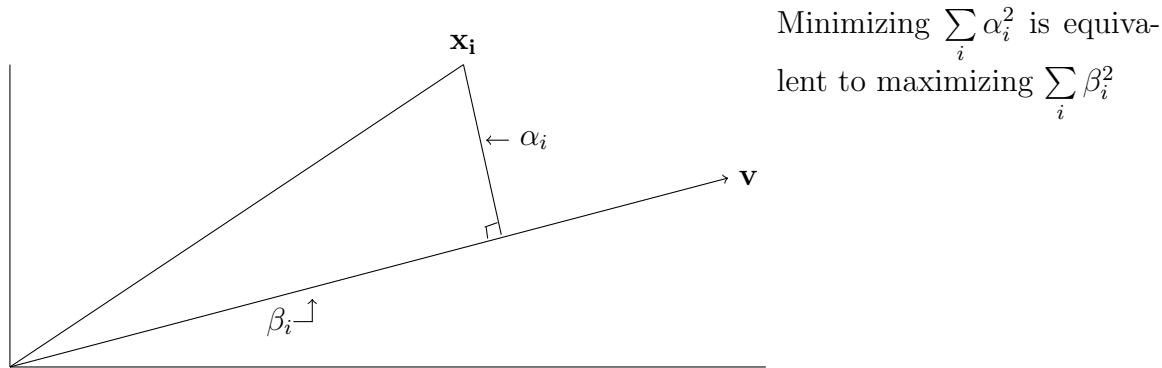


Figure 4.1: The projection of the point $\mathbf{x}_i$ onto the line through the origin in the direction of $\mathbf{v}$.

arbitrary and in a way it is. But the square has many nice mathematical properties. The first of these is the use of Pythagoras theorem to say that minimizing the sum of squared distances is equivalent to maximizing the sum of squared projections.

4.1 Singular Vectors

We now define the *singular vectors* of an $n \times d$ matrix A . Consider the rows of A as n points in a d -dimensional space. Consider the best fit line through the origin. Let $\mathbf{v}$ be a unit vector along this line. The length of the projection of $\mathbf{a}_i$, the i^{th} row of A , onto $\mathbf{v}$ is $|\mathbf{a}_i \cdot \mathbf{v}|$. From this we see that the sum of length squared of the projections is $|\mathbf{A}\mathbf{v}|^2$. The best fit line is the one maximizing $|\mathbf{A}\mathbf{v}|^2$ and hence minimizing the sum of the squared distances of the points to the line.

With this in mind, define the *first singular vector*, $\mathbf{v}_1$, of A , which is a column vector, as the vector defining the best fit line through the origin for the n points in d -space that are the rows of A . Thus

$$\mathbf{v}_1 = \arg \max_{|\mathbf{v}|=1} |\mathbf{A}\mathbf{v}|.$$

Technically $\arg \max$ is a set and if v_1 is not unique we arbitrarily pick one of the vectors and refer to it as "the first singular vector" avoiding the more cumbersome "one of the the vectors achieving the maximum". We adopt this terminology for all uses of $\arg \max$.

The value $\sigma_1(A) = |\mathbf{A}\mathbf{v}_1|$ is called the *first singular value* of A . Note that $\sigma_1^2 = \sum_{i=1}^n (a_i \cdot \mathbf{v}_1)^2$ is the sum of the squares of the projections of the points to the line determined by $\mathbf{v}_1$.

The greedy approach to find the best fit 2-dimensional subspace for a matrix A , takes $\mathbf{v}_1$ as the first basis vector for the 2-dimensional subspace and finds the best 2-dimensional subspace containing $\mathbf{v}_1$. The fact that we are using the sum of squared distances helps. For every 2-dimensional subspace containing $\mathbf{v}_1$, the sum of squared lengths of the pro-

jections onto the subspace equals the sum of squared projections onto $\mathbf{v}_1$ plus the sum of squared projections along a vector perpendicular to $\mathbf{v}_1$ in the subspace. Thus, instead of looking for the best 2-dimensional subspace containing $\mathbf{v}_1$, look for a unit vector $\mathbf{v}_2$ perpendicular to $\mathbf{v}_1$ that maximizes $|A\mathbf{v}|^2$ among all such unit vectors. Using the same greedy strategy to find the best three and higher dimensional subspaces, defines $\mathbf{v}_3, \mathbf{v}_4, \dots$ in a similar manner. This is captured in the following definitions. There is no apriori guarantee that the greedy algorithm gives the best fit. But, in fact, the greedy algorithm does work and yields the best-fit subspaces of every dimension as we will show.

The *second singular vector*, $\mathbf{v}_2$, is defined by the best fit line perpendicular to $\mathbf{v}_1$.

$$\mathbf{v}_2 = \arg \max_{\substack{\mathbf{v} \perp \mathbf{v}_1 \\ |\mathbf{v}|=1}} |A\mathbf{v}|$$

The value $\sigma_2(A) = |A\mathbf{v}_2|$ is called the *second singular value* of A . The *third singular vector* $\mathbf{v}_3$ and *third singular value* are defined similarly by

$$\mathbf{v}_3 = \arg \max_{\substack{\mathbf{v} \perp \mathbf{v}_1, \mathbf{v}_2 \\ |\mathbf{v}|=1}} |A\mathbf{v}|$$

and

$$\sigma_3(A) = |A\mathbf{v}_3|,$$

and so on. The process stops when we have found singular vectors $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r$, singular values $\sigma_1, \sigma_2, \dots, \sigma_r$, and

$$\max_{\substack{\mathbf{v} \perp \mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r \\ |\mathbf{v}|=1}} |A\mathbf{v}| = 0.$$

If instead of finding the $\mathbf{v}_1$ that maximized $|A\mathbf{v}|$ and then the best fit 2-dimensional subspace containing $\mathbf{v}_1$, we had found the best fit 2-dimensional subspace, we might have done better. This is not the case. We now give a simple proof that the greedy algorithm indeed finds the best subspaces of every dimension.

Theorem 4.1 *Let A be an $n \times d$ matrix with singular vectors $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r$. For $1 \leq k \leq r$, let V_k be the subspace spanned by $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_k$. For each k , V_k is the best-fit k -dimensional subspace for A .*

Proof: The statement is obviously true for $k = 1$. For $k = 2$, let W be a best-fit 2-dimensional subspace for A . For any basis $(\mathbf{w}_1, \mathbf{w}_2)$ of W , $|A\mathbf{w}_1|^2 + |A\mathbf{w}_2|^2$ is the sum of squared lengths of the projections of the rows of A onto W . Choose a basis $(\mathbf{w}_1, \mathbf{w}_2)$ of W so that $\mathbf{w}_2$ is perpendicular to $\mathbf{v}_1$. If $\mathbf{v}_1$ is perpendicular to W , any unit vector in W will do as $\mathbf{w}_2$. If not, choose $\mathbf{w}_2$ to be the unit vector in W perpendicular to the projection of $\mathbf{v}_1$ onto W . This makes $\mathbf{w}_2$ perpendicular to $\mathbf{v}_1$. Since $\mathbf{v}_1$ maximizes $|A\mathbf{v}|^2$, it follows that $|A\mathbf{w}_1|^2 \leq |A\mathbf{v}_1|^2$. Since $\mathbf{v}_2$ maximizes $|A\mathbf{v}|^2$ over all $\mathbf{v}$ perpendicular to $\mathbf{v}_1$, $|A\mathbf{w}_2|^2 \leq |A\mathbf{v}_2|^2$. Thus

$$|A\mathbf{w}_1|^2 + |A\mathbf{w}_2|^2 \leq |A\mathbf{v}_1|^2 + |A\mathbf{v}_2|^2.$$

Hence, V_2 is at least as good as W and so is a best-fit 2-dimensional subspace.

For general k , proceed by induction. By the induction hypothesis, V_{k-1} is a best-fit $k-1$ dimensional subspace. Suppose W is a best-fit k -dimensional subspace. Choose a basis $\mathbf{w}_1, \mathbf{w}_2, \dots, \mathbf{w}_k$ of W so that $\mathbf{w}_k$ is perpendicular to $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{k-1}$. Then

$$|A\mathbf{w}_1|^2 + |A\mathbf{w}_2|^2 + \dots + |A\mathbf{w}_k|^2 \leq |A\mathbf{v}_1|^2 + |A\mathbf{v}_2|^2 + \dots + |A\mathbf{v}_{k-1}|^2 + |A\mathbf{w}_k|^2$$

since V_{k-1} is an optimal $k-1$ dimensional subspace. Since $\mathbf{w}_k$ is perpendicular to $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{k-1}$, by the definition of $\mathbf{v}_k$, $|A\mathbf{w}_k|^2 \leq |A\mathbf{v}_k|^2$. Thus

$$|A\mathbf{w}_1|^2 + |A\mathbf{w}_2|^2 + \dots + |A\mathbf{w}_{k-1}|^2 + |A\mathbf{w}_k|^2 \leq |A\mathbf{v}_1|^2 + |A\mathbf{v}_2|^2 + \dots + |A\mathbf{v}_{k-1}|^2 + |A\mathbf{v}_k|^2,$$

proving that V_k is at least as good as W and hence is optimal. ■

Note that the n -vector $A\mathbf{v}_i$ is a list of lengths (with signs) of the projections of the rows of A onto $\mathbf{v}_i$. Think of $|A\mathbf{v}_i| = \sigma_i(A)$ as the “component” of the matrix A along $\mathbf{v}_i$. For this interpretation to make sense, it should be true that adding up the squares of the components of A along each of the $\mathbf{v}_i$ gives the square of the “whole content of the matrix A ”. This is indeed the case and is the matrix analogy of decomposing a vector into its components along orthogonal directions.

Consider one row, say $\mathbf{a}_j$, of A . Since $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r$ span the space of all rows of A , $\mathbf{a}_j \cdot \mathbf{v} = 0$ for all $\mathbf{v}$ perpendicular to $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r$. Thus, for each row $\mathbf{a}_j$, $\sum_{i=1}^r (\mathbf{a}_j \cdot \mathbf{v}_i)^2 = |\mathbf{a}_j|^2$. Summing over all rows j ,

$$\sum_{j=1}^n |\mathbf{a}_j|^2 = \sum_{j=1}^n \sum_{i=1}^r (\mathbf{a}_j \cdot \mathbf{v}_i)^2 = \sum_{i=1}^r \sum_{j=1}^n (\mathbf{a}_j \cdot \mathbf{v}_i)^2 = \sum_{i=1}^r |A\mathbf{v}_i|^2 = \sum_{i=1}^r \sigma_i^2(A).$$

But $\sum_{j=1}^n |\mathbf{a}_j|^2 = \sum_{j=1}^n \sum_{k=1}^d a_{jk}^2$, the sum of squares of all the entries of A . Thus, the sum of squares of the singular values of A is indeed the square of the “whole content of A ”, i.e., the sum of squares of all the entries. There is an important norm associated with this quantity, the Frobenius norm of A , denoted $\|A\|_F$ defined as

$$\|A\|_F = \sqrt{\sum_{j,k} a_{jk}^2}.$$

Lemma 4.2 *For any matrix A , the sum of squares of the singular values equals the square of the Frobenius norm. That is, $\sum \sigma_i^2(A) = \|A\|_F^2$.*

Proof: By the preceding discussion. ■

Let $\mathbf{v}_1$ be the unit vector $\mathbf{v}$ that maximizes $|A\mathbf{v}|$, and for each $i > 1$, let $\mathbf{v}_i$ be the unit vector perpendicular to $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{i-1}$ that maximizes $|A\mathbf{v}|$. The vectors $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r$ are called the *right-singular vectors*. The vectors $A\mathbf{v}_i$ form a fundamental set of vectors and we normalize them to length one by

$$\mathbf{u}_i = \frac{1}{\sigma_i(A)} A\mathbf{v}_i.$$

Later we will show that $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_r$ similarly maximize $|\mathbf{u}^T A|$ when multiplied on the left and are called the *left-singular vectors*.

Clearly, the right-singular vectors are orthogonal by definition. We now show that the left-singular vectors are also orthogonal. The proof is by induction on the rank of the matrix A . The induction hypothesis will state that the left-singular vectors of a rank k matrix A are orthogonal. For A of rank $k + 1$ and left-singular vectors $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_{k+1}$, we argue that $B = A - \sigma_1 \mathbf{u}_1 \mathbf{v}_1^T$ is of rank k and has left-singular vectors $\mathbf{u}_2, \mathbf{u}_3, \dots, \mathbf{u}_{k+1}$, which by the induction hypothesis are orthogonal. Then we show that $\mathbf{u}_1$ is orthogonal to $\mathbf{u}_2, \mathbf{u}_3, \dots, \mathbf{u}_{k+1}$ establishing that the left-singular vectors of A are orthogonal.

Theorem 4.3 *Let A be a rank r matrix with right singular vectors $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r$. The left-singular vectors of A , $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_r$, are orthogonal.*

Proof: The proof is by induction on r . For $r = 1$, there is only one $\mathbf{u}_i$ so the theorem is trivially true. For the inductive part consider the matrix

$$B = A - \sigma_1 \mathbf{u}_1 \mathbf{v}_1^T.$$

The implied algorithm in the definition of singular value decomposition applied to B is identical to the algorithm on A for its second and later singular vectors and singular values. To see this, first observe that $B\mathbf{v}_1 = A\mathbf{v}_1 - \sigma_1 \mathbf{u}_1 \mathbf{v}_1^T \mathbf{v}_1 = 0$. It then follows that the first right-singular vector, call it $\mathbf{z}$, of B will be perpendicular to $\mathbf{v}_1$ since if it had a component $\mathbf{z}_1$ along $\mathbf{v}_1$, then, $\left| B \frac{\mathbf{z} - \mathbf{z}_1}{|\mathbf{z} - \mathbf{z}_1|} \right| = \frac{|B\mathbf{z}|}{|\mathbf{z} - \mathbf{z}_1|} > |B\mathbf{z}|$, contradicting the arg max definition of $\mathbf{z}$. But for any $\mathbf{v}$ perpendicular to $\mathbf{v}_1$, $B\mathbf{v} = A\mathbf{v}$. Thus, the top singular vector of B is indeed a second singular vector of A . Repeating this argument shows that a run of the algorithm on B is the same as a run on A for its second and later singular vectors. This is left as an exercise.

Thus, there is a run of the algorithm that finds that B has right-singular vectors $\mathbf{v}_2, \mathbf{v}_3, \dots, \mathbf{v}_r$ and corresponding left-singular vectors $\mathbf{u}_2, \mathbf{u}_3, \dots, \mathbf{u}_r$. By the induction hypothesis, $\mathbf{u}_2, \mathbf{u}_3, \dots, \mathbf{u}_r$ are orthogonal.

It remains to prove that $\mathbf{u}_1$ is orthogonal to the other $\mathbf{u}_i$. Suppose not, and for some $i \geq 2$, $\mathbf{u}_1^T \mathbf{u}_i \neq 0$. Without loss of generality assume that $\mathbf{u}_1^T \mathbf{u}_i > 0$. The proof is symmetric for the case where $\mathbf{u}_1^T \mathbf{u}_i < 0$. Now, for infinitesimally small $\varepsilon > 0$, the vector

$$A \left(\frac{\mathbf{v}_1 + \varepsilon \mathbf{v}_i}{|\mathbf{v}_1 + \varepsilon \mathbf{v}_i|} \right) = \frac{\sigma_1 \mathbf{u}_1 + \varepsilon \sigma_i \mathbf{u}_i}{\sqrt{1 + \varepsilon^2}}$$

has length at least as large as its component along $\mathbf{u}_1$ which is

$$\mathbf{u}_1^T \left(\frac{\sigma_1 \mathbf{u}_1 + \varepsilon \sigma_i \mathbf{u}_i}{\sqrt{1 + \varepsilon^2}} \right) = (\sigma_1 + \varepsilon \sigma_i \mathbf{u}_1^T \mathbf{u}_i) \left(1 - \frac{\varepsilon^2}{2} + O(\varepsilon^4) \right) = \sigma_1 + \varepsilon \sigma_i \mathbf{u}_1^T \mathbf{u}_i - O(\varepsilon^2) > \sigma_1,$$

a contradiction. Thus, $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_r$ are orthogonal. ■

4.2 Singular Value Decomposition (SVD)

Let A be an $n \times d$ matrix with singular vectors $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r$ and corresponding singular values $\sigma_1, \sigma_2, \dots, \sigma_r$. The left-singular vectors of A are $\mathbf{u}_i = \frac{1}{\sigma_i} A \mathbf{v}_i$ where $\sigma_i \mathbf{u}_i$ is a vector whose coordinates correspond to the projections of the rows of A onto $\mathbf{v}_i$. Each $\sigma_i \mathbf{u}_i \mathbf{v}_i^T$ is a rank one matrix whose columns are weighted versions of the $\sigma_i \mathbf{u}_i$, weighted proportional to the coordinates of $\mathbf{v}_i$.

We will prove that A can be decomposed into a sum the of rank one matrices as

$$A = \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T.$$

We first prove a simple lemma stating that two matrices A and B are identical if $A\mathbf{v} = B\mathbf{v}$ for all $\mathbf{v}$.

Lemma 4.4 *Matrices A and B are identical if and only if for all vectors $\mathbf{v}$, $A\mathbf{v} = B\mathbf{v}$.*

Proof: Clearly, if $A = B$ then $A\mathbf{v} = B\mathbf{v}$ for all $\mathbf{v}$. For the converse, suppose that $A\mathbf{v} = B\mathbf{v}$ for all $\mathbf{v}$. Let $\mathbf{e}_i$ be the vector that is all zeros except for the i^{th} component which has value one. Now $A\mathbf{e}_i$ is the i^{th} column of A and thus $A = B$ if for each i , $A\mathbf{e}_i = B\mathbf{e}_i$. ■

Theorem 4.5 *Let A be an $n \times d$ matrix with right-singular vectors $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r$, left-singular vectors $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_r$, and corresponding singular values $\sigma_1, \sigma_2, \dots, \sigma_r$. Then*

$$A = \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T.$$

Proof: We first show that multiplying A and $\sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T$ by $\mathbf{v}_j$ results in equality.

$$\sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T \mathbf{v}_j = \sigma_j \mathbf{u}_j = A \mathbf{v}_j$$

Since any vector $\mathbf{v}$ can be expressed as a linear combination of the singular vectors plus a vector perpendicular to the $\mathbf{v}_i$, $A\mathbf{v} = \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T \mathbf{v}$ for all $\mathbf{v}$ and by Lemma 4.4,

$$A = \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T. \quad \blacksquare$$

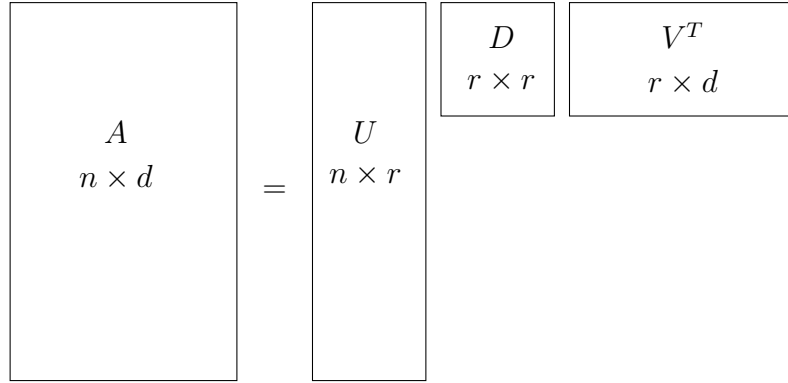


Figure 4.2: The SVD decomposition of an $n \times d$ matrix.

The decomposition $A = \sum_i \sigma_i \mathbf{u}_i \mathbf{v}_i^T$ is called the *singular value decomposition*, *SVD*, of A . In matrix notation $A = UDV^T$ where the columns of U and V consist of the left and right-singular vectors, respectively, and D is a diagonal matrix whose diagonal entries are the singular values of A .

For any matrix A , the sequence of singular values is unique and if the singular values are all distinct, then the sequence of singular vectors is unique also. However, when some set of singular values are equal, the corresponding singular vectors span some subspace. Any set of orthonormal vectors spanning this subspace can be used as the singular vectors.

4.3 Best Rank k Approximations

There are two important matrix norms, the *Frobenius norm* denoted $\|A\|_F$ and the *2-norm* denoted $\|A\|_2$. The 2-norm of the matrix A is given by

$$\max_{\|\mathbf{v}\|=1} |A\mathbf{v}|$$

and thus equals the largest singular value of the matrix.

Let A be an $n \times d$ matrix and think of the rows of A as n points in d -dimensional space. The Frobenius norm of A is the square root of the sum of the squared distance of the points to the origin. The 2-norm is the square root of the sum of squared distances to the origin along the direction that maximizes this quantity. That is, the 2-norm is the square root of the sum of squares of projections onto the first singular vector and thus is given by the first singular value σ_1 .

Let

$$A = \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T$$

be the SVD of A . For $k \in \{1, 2, \dots, r\}$, let

$$A_k = \sum_{i=1}^k \sigma_i \mathbf{u}_i \mathbf{v}_i^T$$

be the sum truncated after k terms. It is clear that A_k has rank k . We now show that A_k is the best rank k approximation to A when the error is measured in either the 2-norm or the Frobenius norm.

Lemma 4.6 *The rows of A_k are the projections of the rows of A onto the subspace V_k spanned by the first k singular vectors of A .*

Proof: Let $\mathbf{a}$ be an arbitrary row vector. Since the $\mathbf{v}_i$ are orthonormal, the projection of the vector $\mathbf{a}$ onto V_k is given by $\sum_{i=1}^k (\mathbf{a} \cdot \mathbf{v}_i) \mathbf{v}_i^T$. Thus, the matrix whose rows are the projections of the rows of A onto V_k is given by $\sum_{i=1}^k A \mathbf{v}_i \mathbf{v}_i^T$. This last expression simplifies to

$$\sum_{i=1}^k A \mathbf{v}_i \mathbf{v}_i^T = \sum_{i=1}^k \sigma_i \mathbf{u}_i \mathbf{v}_i^T = A_k.$$

■

The matrix A_k is the best rank k approximation to A in both the Frobenius and the 2-norm. First we show that the matrix A_k is the best rank k approximation to A in the Frobenius norm.

Theorem 4.7 *For any matrix B of rank at most k*

$$\|A - A_k\|_F \leq \|A - B\|_F$$

Proof: Let B minimize $\|A - B\|_F^2$ among all rank k or less matrices. Let V be the space spanned by the rows of B . The dimension of V is at most k . Since B minimizes $\|A - B\|_F^2$, it must be that each row of B is the projection of the corresponding row of A onto V , otherwise replacing the row of B with the projection of the corresponding row of A onto V does not change V and hence the rank of B but would reduce $\|A - B\|_F^2$. Since now each row of B is the projection of the corresponding row of A , it follows that $\|A - B\|_F^2$ is the sum of squared distances of rows of A to V . Since A_k minimizes the sum of squared distance of rows of A to any k -dimensional subspace, from Theorem 4.1, it follows that $\|A - A_k\|_F \leq \|A - B\|_F$.

■

Next we tackle the 2-norm. We first show that the square of the 2-norm of $A - A_k$ is the square of the $(k + 1)^{st}$ singular value of A ,

Lemma 4.8 $\|A - A_k\|_2^2 = \sigma_{k+1}^2$.

Proof: Let $A = \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T$ be the singular value decomposition of A . Then $A_k = \sum_{i=1}^k \sigma_i \mathbf{u}_i \mathbf{v}_i^T$ and $A - A_k = \sum_{i=k+1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T$. Let $\mathbf{v}$ be the top singular vector of $A - A_k$. Express $\mathbf{v}$ as a linear combination of $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r$. That is, write $\mathbf{v} = \sum_{i=1}^r \alpha_i \mathbf{v}_i$. Then

$$\begin{aligned} |(A - A_k)\mathbf{v}| &= \left| \sum_{i=k+1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T \sum_{j=1}^r \alpha_j \mathbf{v}_j \right| = \left| \sum_{i=k+1}^r \alpha_i \sigma_i \mathbf{u}_i \mathbf{v}_i^T \mathbf{v}_i \right| \\ &= \left| \sum_{i=k+1}^r \alpha_i \sigma_i \mathbf{u}_i \right| = \sqrt{\sum_{i=k+1}^r \alpha_i^2 \sigma_i^2}. \end{aligned}$$

The $\mathbf{v}$ maximizing this last quantity, subject to the constraint that $|\mathbf{v}|^2 = \sum_{i=1}^r \alpha_i^2 = 1$, occurs when $\alpha_{k+1} = 1$ and the rest of the α_i are 0. Thus, $\|A - A_k\|_2^2 = \sigma_{k+1}^2$ proving the lemma. ■

Finally, we prove that A_k is the best rank k , 2-norm approximation to A .

Theorem 4.9 *Let A be an $n \times d$ matrix. For any matrix B of rank at most k*

$$\|A - A_k\|_2 \leq \|A - B\|_2.$$

Proof: If A is of rank k or less, the theorem is obviously true since $\|A - A_k\|_2 = 0$. Assume that A is of rank greater than k . By Lemma 4.8, $\|A - A_k\|_2^2 = \sigma_{k+1}^2$. Suppose there is some matrix B of rank at most k such that B is a better 2-norm approximation to A than A_k . That is, $\|A - B\|_2 < \sigma_{k+1}$. The null space of B , the set of vectors $\mathbf{v}$ such that $B\mathbf{v} = 0$, has dimension at least $d - k$. Let $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{k+1}$ be the first $k + 1$ singular vectors of A . By a dimension argument, it follows that there exists a $\mathbf{z} \neq 0$ in

$$\text{Null}(B) \cap \text{Span}\{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{k+1}\}.$$

Let $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_{d-k}$ be $d - k$ independent vectors in $\text{Null}(B)$. Now, $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_{d-k}, \mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{k+1}$ are $d + 1$ vectors in d space and thus are linearly dependent. Let $\alpha_1, \alpha_2, \dots, \alpha_{d-k}$ and $\beta_1, \beta_2, \dots, \beta_k$ be such that $\sum_{i=1}^{d-k} \alpha_i \mathbf{u}_i = \sum_{j=1}^k \beta_j \mathbf{v}_j$. Let $\mathbf{z} = \sum_{i=1}^{d-k} \alpha_i \mathbf{u}_i$. Scale $\mathbf{z}$ so that $|\mathbf{z}| = 1$. We now show that for this vector $\mathbf{z}$, which lies in the space of the first $k + 1$ singular vectors of A , that $(A - B)\mathbf{z} \geq \sigma_{k+1}$. Hence the 2-norm of $A - B$ is at least σ_{k+1} , contradicting the assumption that $\|A - B\|_2 < \sigma_{k+1}$. First

$$\|A - B\|_2^2 \geq |(A - B)\mathbf{z}|^2.$$

Since $B\mathbf{z} = 0$,

$$\|A - B\|_2^2 \geq |A\mathbf{z}|^2.$$

Since $\mathbf{z}$ is in the Span $\{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{k+1}\}$

$$|A\mathbf{z}|^2 = \left| \sum_{i=1}^n \sigma_i \mathbf{u}_i \mathbf{v}_i^T \mathbf{z} \right|^2 = \sum_{i=1}^n \sigma_i^2 (\mathbf{v}_i^T \mathbf{z})^2 = \sum_{i=1}^{k+1} \sigma_i^2 (\mathbf{v}_i^T \mathbf{z})^2 \geq \sigma_{k+1}^2 \sum_{i=1}^{k+1} (\mathbf{v}_i^T \mathbf{z})^2 = \sigma_{k+1}^2.$$

It follows that

$$\|A - B\|_2^2 \geq \sigma_{k+1}^2$$

contradicting the assumption that $\|A - B\|_2 < \sigma_{k+1}$. This proves the theorem. ■

4.4 Power Method for Computing the Singular Value Decomposition

Computing the singular value decomposition is an important branch of numerical analysis in which there have been many sophisticated developments over a long period of time. Here we present an “in-principle” method to establish that the approximate SVD of a matrix A can be computed in polynomial time. The reader is referred to numerical analysis texts for more details. The method we present, called the *power method*, is simple and is in fact the conceptual starting point for many algorithms. Let A be a matrix whose SVD is $\sum_i \sigma_i \mathbf{u}_i \mathbf{v}_i^T$. We wish to work with a matrix that is square and symmetric. By direct multiplication, since $\mathbf{u}_i^T \mathbf{u}_j$ is the dot product of the two vectors and is zero unless $i = j$

$$\begin{aligned} B &= A^T A = \left(\sum_i \sigma_i \mathbf{v}_i \mathbf{u}_i^T \right) \left(\sum_j \sigma_j \mathbf{u}_j \mathbf{v}_j^T \right) \\ &= \sum_{i,j} \sigma_i \sigma_j \mathbf{v}_i (\mathbf{u}_i^T \cdot \mathbf{u}_j) \mathbf{v}_j^T = \sum_i \sigma_i^2 \mathbf{v}_i \mathbf{v}_i^T, \end{aligned}$$

The matrix B is square and symmetric, and has the same left and right-singular vectors. If A is itself square and symmetric, it will have the same right and left-singular vectors, namely $A = \sum_i \sigma_i \mathbf{v}_i \mathbf{v}_i^T$ and computing B is unnecessary.

Now consider computing B^2 .

$$B^2 = \left(\sum_i \sigma_i^2 \mathbf{v}_i \mathbf{v}_i^T \right) \left(\sum_j \sigma_j^2 \mathbf{v}_j \mathbf{v}_j^T \right) = \sum_{ij} \sigma_i^2 \sigma_j^2 \mathbf{v}_i (\mathbf{v}_i^T \mathbf{v}_j) \mathbf{v}_j^T$$

When $i \neq j$, the dot product $\mathbf{v}_i^T \mathbf{v}_j$ equals 0 by orthogonality. However the “outer product” $\mathbf{v}_i \mathbf{v}_j^T$ is a matrix and is not zero even for $i \neq j$. Thus, $B^2 = \sum_{i=1}^r \sigma_i^4 \mathbf{v}_i \mathbf{v}_i^T$. In computing the k^{th} power of B , again all the cross product terms are zero and

$$B^k = \sum_{i=1}^r \sigma_i^{2k} \mathbf{v}_i \mathbf{v}_i^T.$$

If $\sigma_1 > \sigma_2$, then

$$\frac{1}{\sigma_1^{2k}} B^k \rightarrow \mathbf{v}_1 \mathbf{v}_1^T.$$

We do not know σ_1 beforehand and cannot find this limit. If we take B^k and divide by $\|B^k\|_F$ so that the Frobenius norm is normalized to one, the matrix will converge to the rank one matrix $\mathbf{v}_1 \mathbf{v}_1^T$ from which $\mathbf{v}_1$ may be computed by normalizing the first column to be a unit vector.

The difficulty with the above method is that A may be a very large, sparse matrix, say a $10^8 \times 10^8$ matrix with 10^9 nonzero entries. In this case, B will be a $10^8 \times 10^8$ matrix, but will not be sparse. We might not even be able to store B , let alone compute the product B^2 . Even if A is moderate in size, computing matrix products is costly in time. Thus, we need a more efficient method.

Instead of computing B^k , select a random vector $\mathbf{x}$ and compute the product $B^k \mathbf{x}$. The vector $\mathbf{x}$ can be expressed in terms of the singular vectors of B augmented to a full basis as $\mathbf{x} = \sum c_i \mathbf{v}_i$. Then

$$B^k \mathbf{x} \approx (\sigma_1^{2k} \mathbf{v}_1 \mathbf{v}_1^T) \left(\sum_{i=1}^n c_i \mathbf{v}_i \right) = \sigma_1^{2k} c_1 \mathbf{v}_1$$

Normalizing the resulting vector yields $\mathbf{v}_1$, the first singular vector of A . The way $B^k \mathbf{x}$ is computed is by a series of matrix vector products, instead of matrix products. $B \mathbf{x} = A(A \mathbf{x})$ and $B^k \mathbf{x} = (A^T A B^{k-1} \mathbf{x})$. Thus, we perform $2k$ vector times sparse matrix multiplications.

An issue occurs if there is no significant gap between the first and second singular values of a matrix. Then the power method will not quickly converge to the first left-singular vector. There are ways of overcoming this. Theorem 4.11 states that even with ties, the power method converges to some vector in the span of those singular vectors corresponding to the “nearly highest” singular values.

We start with a technical lemma needed in the proof of the theorem.

Lemma 4.10 *Let $(x_1, x_2, \dots, x_d)$ be a unit d -dimensional vector picked at random from the set $\{\mathbf{x} : |\mathbf{x}| = 1\}$. The probability that $|x_1| \geq \frac{1}{20\sqrt{d}}$ is at least $9/10$.*

Proof: We first show that for a vector $\mathbf{v}$ picked at random with $|\mathbf{v}| \leq 1$, the probability that $v_1 \geq \frac{1}{20\sqrt{d}}$ is at least $9/10$. Then we let $\mathbf{x} = \mathbf{v}/|\mathbf{v}|$. This can only increase the value of v_1 , so the result follows.

Let $\alpha = \frac{1}{20\sqrt{d}}$. The probability that $|v_1| \geq \alpha$ equals one minus the probability that $|v_1| \leq \alpha$. The probability that $|v_1| \leq \alpha$ is equal to the fraction of the volume of the unit

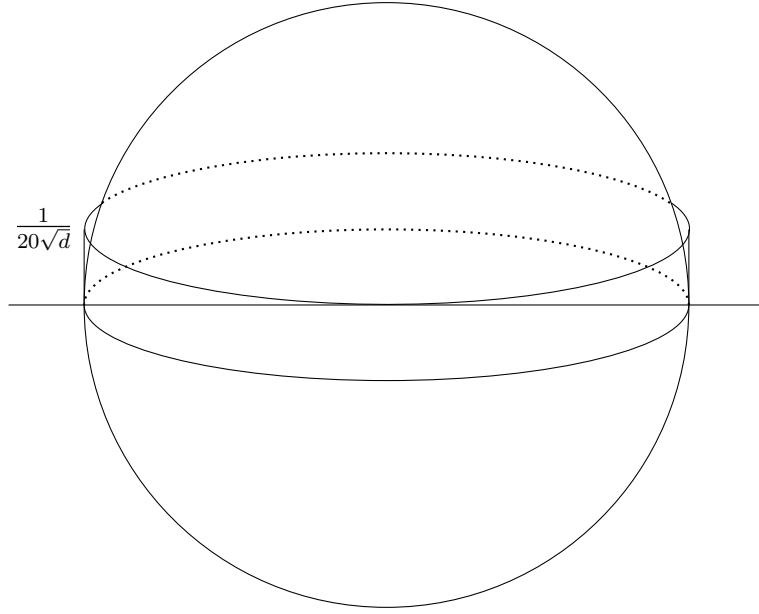


Figure 4.3: The volume of the cylinder of height $\frac{1}{20\sqrt{d}}$ is an upper bound on the volume of the hemisphere below $x_1 = \frac{1}{20\sqrt{d}}$.

sphere with $|v_1| \leq \alpha$. To get an upper bound on the volume of the sphere with $|v_1| \leq \alpha$, consider twice the volume of the unit radius cylinder of height α . The volume of the portion of the sphere with $|v_1| \leq \alpha$ is less than or equal to $2\alpha V(d-1)$ (recall notation from Chapter 2) and

$$\text{Prob}(|v_1| \leq \alpha) \leq \frac{2\alpha V(d-1)}{V(d)}.$$

Now the volume of the unit radius sphere is at least twice the volume of the cylinder of height $\frac{1}{\sqrt{d-1}}$ and radius $\sqrt{1 - \frac{1}{d-1}}$ or

$$V(d) \geq \frac{2}{\sqrt{d-1}} V(d-1) \left(1 - \frac{1}{d-1}\right)^{\frac{d-1}{2}}.$$

Using $(1-x)^a \geq 1-ax$

$$V(d) \geq \frac{2}{\sqrt{d-1}} V(d-1) \left(1 - \frac{d-1}{2} \frac{1}{d-1}\right) \geq \frac{V(d-1)}{\sqrt{d-1}}$$

and

$$\text{Prob}(|v_1| \leq \alpha) \leq \frac{2\alpha V(d-1)}{\frac{1}{\sqrt{d-1}} V(d-1)} \leq \frac{\sqrt{d-1}}{10\sqrt{d}} \leq \frac{1}{10}.$$

Thus, the probability that $v_1 \geq \frac{1}{20\sqrt{d}}$ is at least 9/10. ■

Theorem 4.11 Let A be an $n \times d$ matrix and $\mathbf{x}$ a random unit length vector. Let V be the space spanned by the left-singular vectors of A corresponding to singular values greater than $(1 - \varepsilon) \sigma_1$. Let k be $\Omega\left(\frac{\ln(n/\varepsilon)}{\varepsilon}\right)$. Let $\mathbf{w}$ be unit vector after k iterations of the power method, namely,

$$\mathbf{w} = \frac{(AA^T)^k \mathbf{x}}{\|(AA^T)^k \mathbf{x}\|}.$$

The probability that $\mathbf{w}$ has a component of at least ε perpendicular to V is at most $1/10$.

Proof: Let

$$A = \sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T$$

be the SVD of A . If the rank of A is less than n , then complete $\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_r\}$ into a basis $\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_n\}$ of n -space. Write $\mathbf{x}$ in the basis of the $\mathbf{u}_i$'s as

$$\mathbf{x} = \sum_{i=1}^n c_i \mathbf{u}_i.$$

Since $(AA^T)^k = \sum_{i=1}^n \sigma_i^{2k} \mathbf{u}_i \mathbf{u}_i^T$, it follows that $(AA^T)^k \mathbf{x} = \sum_{i=1}^n \sigma_i^{2k} c_i \mathbf{u}_i$. For a random unit length vector $\mathbf{x}$ picked independent of A , the $\mathbf{u}_i$ are fixed vectors and picking $\mathbf{x}$ at random is equivalent to picking random c_i . From Lemma 4.10, $|c_1| \geq \frac{1}{20\sqrt{n}}$ with probability at least $9/10$.

Suppose that $\sigma_1, \sigma_2, \dots, \sigma_m$ are the singular values of A that are greater than or equal to $(1 - \varepsilon) \sigma_1$ and that $\sigma_{m+1}, \dots, \sigma_n$ are the singular values that are less than $(1 - \varepsilon) \sigma_1$. Now

$$\|(AA^T)^k \mathbf{x}\|^2 = \left\| \sum_{i=1}^n \sigma_i^{2k} c_i \mathbf{u}_i \right\|^2 = \sum_{i=1}^n \sigma_i^{4k} c_i^2 \geq \sigma_1^{4k} c_1^2 \geq \frac{1}{400d} \sigma_1^{4k},$$

with probability at least $9/10$. Here we used the fact that a sum of positive quantities is at least as large as its first element and the first element is greater than or equal to $\frac{1}{400d} \sigma_1^{4k}$ with probability at least $9/10$. If we did not choose $\mathbf{x}$ at random in the beginning and use Lemma 4.10, c_1 could well have been zero and this argument would not work.

The component of $\|(AA^T)^k \mathbf{x}\|^2$ perpendicular to the space V is

$$\sum_{i=m+1}^n \sigma_i^{4k} c_i^2 \leq (1 - \varepsilon)^{4k} \sigma_1^{4k} \sum_{i=m+1}^n c_i^2 \leq (1 - \varepsilon)^{4k} \sigma_1^{4k}$$

since $\sum_{i=1}^n c_i^2 = \|\mathbf{x}\|^2 = 1$. Thus, the component of $\mathbf{w}$ perpendicular to V is at most

$$\frac{(1 - \varepsilon)^{2k} \sigma_1^{2k}}{\frac{1}{20\sqrt{n}} \sigma_1^{2k}} = O(\sqrt{n}(1 - \varepsilon)^{2k}) = O(\sqrt{n}e^{-2\varepsilon k}) = O(\sqrt{n}e^{-\Omega(\ln(n/\varepsilon))}) \leq \varepsilon$$

for a suitable choice of constant in the definition of k . ■

4.5 Applications of Singular Value Decomposition

4.5.1 Principal Component Analysis

The traditional use of SVD is in Principal Component Analysis (PCA). PCA is illustrated by a customer-product data problem where there are n customers buying d products. Let matrix A with elements a_{ij} represent the probability of customer i purchasing product j . One hypothesizes that there are only k underlying basic factors like age, income, family size, etc. that determine a customer's purchase behavior. An individual customer's behavior is determined by some weighted combination of these underlying factors. That is, a customer's purchase behavior can be characterized by a k -dimensional vector where k is much smaller than n or d . The components of the vector are weights for each of the basic factors. Associated with each basic factor is a vector of probabilities, each component of which is the probability of purchasing a given product by someone whose behavior depends only on that factor. More abstractly, A is an $n \times d$ matrix that can be expressed as the product of two matrices U and V where U is an $n \times k$ matrix expressing the factor weights for each customer and V is a $k \times d$ matrix expressing the purchase probabilities of products that correspond to that factor. Finding the best rank k approximation A_k by SVD gives such a U and V . One twist is that A may not be exactly equal to UV , but close to it since there may be noise or random perturbations in which case $A - UV$ is treated as noise.

In the above setting, A was available fully and we wished to find U and V to identify the basic factors. If n and d are very large, on the order of thousands or even millions, there is probably little one could do to estimate or even store A . In this setting, we may assume that we are given just a few elements of A and wish to estimate A . If A was an arbitrary matrix of size $n \times d$, this would require $\Omega(nd)$ pieces of information and cannot be done with a few entries. But again hypothesize that A was a small rank matrix with added noise. If now we also assume that the given entries are randomly drawn according to some known distribution, then there is a possibility that SVD can be used to estimate the whole of A . This area is called collaborative filtering and one of its uses is to target an ad to a customer based on one or two purchases. We do not describe it here.

4.5.2 Clustering a Mixture of Spherical Gaussians

Clustering, is the task of partitioning a set of points in d -space into k subsets or clusters where each cluster consists of "nearby" points. Different definitions of the goodness of a clustering lead to different solutions. Clustering is an important area which we will study in detail in Chapter ???. Here we will see how to solve a particular clustering problem using singular value decomposition.

In general, a solution to any clustering problem comes up with k *cluster centers* that define the k clusters. A cluster is the set of data points that are closest to a particular cluster center. Hence the Voronoi cells of the cluster centers determine the clusters. Using

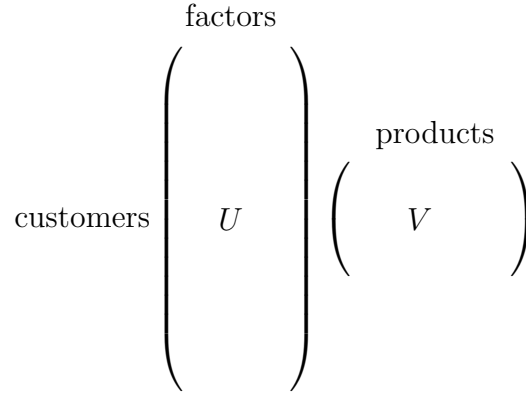


Figure 4.4: Customer-product data

this observation, it is relatively easy to cluster points in two or three dimensions. However, clustering is not so easy in higher dimensions. Many problems have high-dimensional data and clustering problems are no exception.

Clustering problems tend to be NP-hard, so there are no polynomial time algorithms to solve them. One way around this is to assume stochastic models of input data and devise algorithms to cluster data generated by such models. Mixture models are a very important class of stochastic models. A mixture is a probability density or distribution that is the weighted sum of simple component probability densities. It is of the form $w_1p_1 + w_2p_2 + \dots + w_kp_k$ where $p_1, p_2, \dots, p_k$ are the basic probability densities and $w_1, w_2, \dots, w_k$ are positive real numbers called weights that add up to one. Clearly, $w_1p_1 + w_2p_2 + \dots + w_kp_k$ is a probability density, it integrates to one.

The *model fitting problem* is to fit a mixture of k basic densities to n independent, identically distributed samples, each sample drawn according to the same mixture distribution. The class of basic densities is known, but various parameters such as their mean and the component weights of the mixture are not. Here, we deal with the case where the basic densities are all spherical Gaussians. The samples are generated by picking an integer i from the set $\{1, 2, \dots, k\}$ with probabilities $w_1, w_2, \dots, w_k$ and then picking a sample according to p_i . The process is repeated n times. This process generates n samples according to the mixture model where the set of samples is naturally partitioned into k sets, each set corresponding to one p_i .

The model-fitting problem consists of two sub problems. The first sub problem is to cluster the sample into k subsets where each subset corresponds to one component density. The second sub problem is to fit a distribution to each subset. We discuss only the clustering problem here. The problem of fitting a single Gaussian to a set of data points was discussed in Section 2.6 of Chapter 2.

If the component Gaussians in the mixture have their centers very close together, then the clustering problem is unresolvable. In the limiting case where a pair of component densities are the same, there is no way to distinguish between them. What condition on the inter-center separation will guarantee unambiguous clustering? First, by looking at 1-dimensional examples, it is clear that this separation should be measured in units of the standard deviation, since the density is a function of the number of standard deviation from the mean. In one dimension, if two Gaussians have inter-center separation at least three times the maximum of their standard deviations, then they hardly overlap. What is the analog of this statement in higher dimensions?

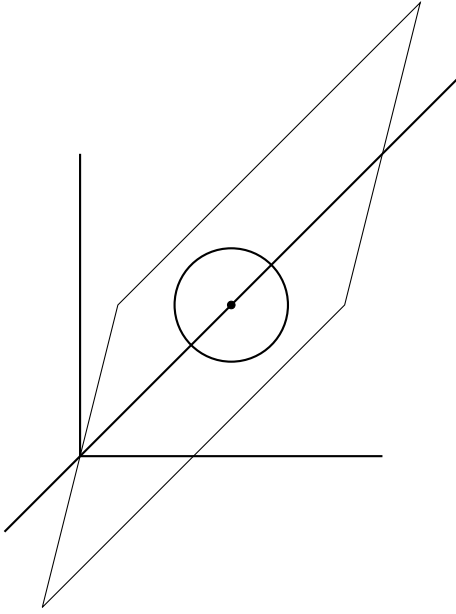
For a d -dimensional spherical Gaussian, with standard deviation σ in each direction², it is easy to see that the expected distance squared from the center is $d\sigma^2$. Define the radius r of the Gaussian to be the square root of the average distance squared from the center; so r is $\sqrt{d}\sigma$. If the inter-center separation between two spherical Gaussians, both of radius r , is at least $2r = 2\sqrt{d}\sigma$, then it is easy to see that the densities hardly overlap. But this separation requirement grows with d . For problems with large d , this would impose a separation requirement not met in practice. The main aim is to answer affirmatively the question: Can we show an analog of the 1-dimensional statement for large d ? In a mixture of k spherical Gaussians in d -space, if the centers of each pair of component Gaussians are $\Omega(1)$ standard deviations apart, then we can separate the sample points into the k components for $k \in O(1)$.

The central idea is the following. Suppose we can find the subspace spanned by the k centers and project the sample points to this subspace. The projection of a spherical Gaussian with standard deviation σ remains a spherical Gaussian with standard deviation σ (Lemma 4.12). In the projection, the inter-center separation remains the same. So in the projection, the Gaussians are distinct provided the inter-center separation in the whole space is $\Omega(\sqrt{k}\sigma)$ which is a lot smaller than the $\Omega(\sqrt{d}\sigma)$ for $k \ll d$. Interestingly the subspace spanned by the k -centers is essentially the best-fit k -dimensional subspace that can be found by singular value decomposition.

Lemma 4.12 *Suppose p is a d -dimensional spherical Gaussian with center μ and standard deviation σ . The density of p projected onto a k -dimensional subspace V is a spherical Gaussian with the same standard deviation.*

Proof: Rotate the coordinate system so V is spanned by the first k coordinate vectors. The Gaussian remains spherical with standard deviation σ although the coordinates of its center have changed. For a point $\mathbf{x} = (x_1, x_2, \dots, x_d)$, we will use the notation $\mathbf{x}' = (x_1, x_2, \dots, x_k)$ and $\mathbf{x}'' = (x_{k+1}, x_{k+2}, \dots, x_n)$. The density of the projected Gaussian at

²Since a spherical Gaussian has the same standard deviation in every direction, we call it the standard deviation of the Gaussian.



1. The best fit 1-dimension subspace to a spherical Gaussian is the line through its center and the origin.
2. Any k -dimensional subspace containing the line is a best fit k -dimensional subspace for the Gaussian.
3. The best fit k -dimensional subspace for k spherical Gaussians is the subspace containing their centers.

Figure 4.5: Best fit subspace to a spherical Gaussian.

the point $(x_1, x_2, \dots, x_k)$ is

$$ce^{-\frac{|\mathbf{x}' - \boldsymbol{\mu}'|^2}{2\sigma^2}} \int_{\mathbf{x}''} e^{-\frac{|\mathbf{x}'' - \boldsymbol{\mu}''|^2}{2\sigma^2}} d\mathbf{x}'' = c'e^{-\frac{|\mathbf{x}' - \boldsymbol{\mu}'|^2}{2\sigma^2}}.$$

This clearly implies the lemma. ■

We now show that the top k singular vectors produced by the SVD span the space of the k centers. First, we extend the notion of best fit to probability distributions. Then we show that for a single spherical Gaussian whose center is not the origin, the best fit 1-dimensional subspace is the line through the center of the Gaussian and the origin. Next, we show that the best fit k -dimensional subspace for a single Gaussian whose center is not the origin is any k -dimensional subspace containing the line through the Gaussian's center and the origin. Finally, for k spherical Gaussians, the best fit k -dimensional subspace is the subspace containing their centers. Thus, the SVD finds the subspace that contains the centers.

Recall that for a set of points, the best-fit line is the line passing through the origin that minimizes the sum of squared distances to the points. We extend this definition to probability densities instead of a set of points.

Definition 4.1 *If p is a probability density in d space, the best fit line for p is the line l passing through the origin that minimizes the expected squared perpendicular distance to*

the line, namely,

$$\int \text{dist}(\mathbf{x}, l)^2 p(\mathbf{x}) d\mathbf{x}.$$

■

Recall that a k -dimensional subspace is the best-fit subspace if the sum of squared distances to it is minimized or equivalently, the sum of squared lengths of projections onto it is maximized. This was defined for a set of points, but again it can be extended to a density as above.

Definition 4.2 *If p is a probability density in d -space and V is a subspace, then the expected squared perpendicular distance of V to p , denoted $f(V, p)$, is given by*

$$f(V, p) = \int (\text{dist}(\mathbf{x}, V))^2 p(\mathbf{x}) d\mathbf{x},$$

where $\text{dist}(\mathbf{x}, V)$ denotes the perpendicular distance from the point $\mathbf{x}$ to the subspace V .

■

For the uniform density on the unit circle centered at the origin, it is easy to see that any line passing through the origin is a best fit line for the probability distribution. Our next lemma shows that the best fit line for a spherical Gaussian centered at $\boldsymbol{\mu} \neq 0$ is the line passing through $\boldsymbol{\mu}$ and the origin.

Lemma 4.13 *Let the probability density p be a spherical Gaussian with center $\boldsymbol{\mu} \neq 0$. The best fit 1-dimensional subspace is the line passing through $\boldsymbol{\mu}$ and the origin.*

Proof: For a randomly chosen $\mathbf{x}$ (according to p) and a fixed unit length vector $\mathbf{v}$,

$$\begin{aligned} E[(\mathbf{v}^T \mathbf{x})^2] &= E[(\mathbf{v}^T (\mathbf{x} - \boldsymbol{\mu}) + \mathbf{v}^T \boldsymbol{\mu})^2] \\ &= E[(\mathbf{v}^T (\mathbf{x} - \boldsymbol{\mu}))^2 + 2(\mathbf{v}^T \boldsymbol{\mu})(\mathbf{v}^T (\mathbf{x} - \boldsymbol{\mu})) + (\mathbf{v}^T \boldsymbol{\mu})^2] \\ &= E[(\mathbf{v}^T (\mathbf{x} - \boldsymbol{\mu}))^2] + 2(\mathbf{v}^T \boldsymbol{\mu}) E[\mathbf{v}^T (\mathbf{x} - \boldsymbol{\mu})] + (\mathbf{v}^T \boldsymbol{\mu})^2 \\ &= E[(\mathbf{v}^T (\mathbf{x} - \boldsymbol{\mu}))^2] + (\mathbf{v}^T \boldsymbol{\mu})^2 \\ &= \sigma^2 + (\mathbf{v}^T \boldsymbol{\mu})^2 \end{aligned}$$

since $E[(\mathbf{v}^T (\mathbf{x} - \boldsymbol{\mu}))^2]$ is the variance in the direction $\mathbf{v}$ and $E[\mathbf{v}^T (\mathbf{x} - \boldsymbol{\mu})] = 0$. The lemma follows from the fact that the best fit line $\mathbf{v}$ is the one that maximizes $(\mathbf{v}^T \boldsymbol{\mu})^2$ which is maximized when $\mathbf{v}$ is aligned with the center $\boldsymbol{\mu}$.

■

Lemma 4.14 *For a spherical Gaussian with center $\boldsymbol{\mu}$, a k -dimensional subspace is a best fit subspace if and only if it contains $\boldsymbol{\mu}$.*

Proof: By symmetry, every k -dimensional subspace through $\boldsymbol{\mu}$ has the same sum of distances squared to the density. Now by the SVD procedure, we know that the best-fit k -dimensional subspace contains the best fit line, i.e., contains $\boldsymbol{\mu}$. Thus, the lemma follows. ■

This leads to the following theorem.

Theorem 4.15 *If p is a mixture of k spherical Gaussians whose centers span a k -dimensional subspace, then the best fit k -dimensional subspace is the one containing the centers.*

Proof: Let p be the mixture $w_1p_1 + w_2p_2 + \cdots + w_kp_k$. Let V be any subspace of dimension k or less. Then, the expected squared perpendicular distance of V to p is

$$\begin{aligned} f(V, p) &= \int \text{dist}^2(\mathbf{x}, V) p(\mathbf{x}) d\mathbf{x} \\ &= \sum_{i=1}^k w_i \int \text{dist}^2(\mathbf{x}, V) p_i(\mathbf{x}) d\mathbf{x} \\ &\geq \sum_{i=1}^k w_i (\text{distance squared of } p_i \text{ to its best fit } k\text{-dimensional subspace}). \end{aligned}$$

Choose V to be the space spanned by the centers of the densities p_i . By Lemma ?? the last inequality becomes an equality proving the theorem. ■

For an infinite set of points drawn according to the mixture, the k -dimensional SVD subspace gives exactly the space of the centers. In reality, we have only a large number of samples drawn according to the mixture. However, it is intuitively clear that as the number of samples increases, the set of sample points approximates the probability density and so the SVD subspace of the sample is close to the space spanned by the centers. The details of how close it gets as a function of the number of samples are technical and we do not carry this out here.

4.5.3 Spectral Decomposition

Let B be a square matrix. If the vector $\mathbf{x}$ and scalar λ are such that $B\mathbf{x} = \lambda\mathbf{x}$, then $\mathbf{x}$ is an *eigenvector* of the matrix B and λ is the corresponding *eigenvalue*. We present here a spectral decomposition theorem for the special case where B is of the form $B = AA^T$ for some possibly rectangular matrix A . If A is a real valued matrix, then B is symmetric and positive definite. That is, $\mathbf{x}^T B \mathbf{x} > 0$ for all nonzero vectors $\mathbf{x}$. The spectral decomposition theorem holds more generally and the interested reader should consult a linear algebra book.

Theorem 4.16 (Spectral Decomposition) *If $B = AA^T$ then $B = \sum_i \sigma_i^2 \mathbf{u}_i \mathbf{u}_i^T$ where $A = \sum_i \sigma_i \mathbf{u}_i \mathbf{v}_i^T$ is the singular valued decomposition of A .*

Proof:

$$\begin{aligned} B = AA^T &= \left(\sum_i \sigma_i \mathbf{u}_i \mathbf{v}_i^T \right) \left(\sum_j \sigma_j \mathbf{u}_j \mathbf{v}_j^T \right)^T \\ &= \sum_i \sum_j \sigma_i \sigma_j \mathbf{u}_i \mathbf{v}_i^T \mathbf{v}_j \mathbf{u}_j^T \\ &= \sum_i \sigma_i^2 \mathbf{u}_i \mathbf{u}_i^T. \end{aligned}$$

■

When the σ_i are all distinct, the $\mathbf{u}_i$ are the eigenvectors of B and the σ_i^2 are the corresponding eigenvalues. If the σ_i are not distinct, then any vector that is a linear combination of those $\mathbf{u}_i$ with the same eigenvalue is an eigenvector of B .

4.5.4 Singular Vectors and Ranking Documents

An important task for a document collection is to rank the documents according to their intrinsic relevance to the collection. A good candidate is a document's projection onto the best-fit direction for the collection of term-document vectors, namely the top left-singular vector of the term-document matrix. An intuitive reason for this is that this direction has the maximum sum of squared projections of the collection and so can be thought of as a synthetic term-document vector best representing the document collection.

Ranking in order of the projection of each document's term vector along the best fit direction has a nice interpretation in terms of the power method. For this, we consider a different example, that of the web with hypertext links. The World Wide Web can be represented by a directed graph whose nodes correspond to web pages and directed edges to hypertext links between pages. Some web pages, called *authorities*, are the most prominent sources for information on a given topic. Other pages called *hubs*, are ones that identify the authorities on a topic. Authority pages are pointed to by many hub pages and hub pages point to many authorities. One is led to what seems like a circular definition: a hub is a page that points to many authorities and an authority is a page that is pointed to by many hubs.

One would like to assign hub weights and authority weights to each node of the web. If there are n nodes, the hub weights form a n -dimensional vector $\mathbf{u}$ and the authority weights form a n -dimensional vector $\mathbf{v}$. Suppose A is the adjacency matrix representing the directed graph. Here a_{ij} is 1 if there is a hypertext link from page i to page j and 0

otherwise. Given hub vector $\mathbf{u}$, the authority vector $\mathbf{v}$ could be computed by the formula

$$v_j = \sum_{i=1}^d u_i a_{ij}$$

since the right hand side is the sum of the hub weights of all the nodes that point to node j . In matrix terms,

$$\mathbf{v} = A^T \mathbf{u}.$$

Similarly, given an authority vector $\mathbf{v}$, the hub vector $\mathbf{u}$ could be computed by $\mathbf{u} = A\mathbf{v}$. Of course, at the start, we have neither vector. But the above discussion suggests a power iteration. Start with any $\mathbf{v}$. Set $\mathbf{u} = A\mathbf{v}$; then set $\mathbf{v} = A^T \mathbf{u}$ and repeat the process. We know from the power method that this converges to the left and right-singular vectors. So after sufficiently many iterations, we may use the left vector $\mathbf{u}$ as hub weights vector and project each column of A onto this direction and rank columns (authorities) in order of this projection. But the projections just form the vector $A^T \mathbf{u}$ which equals $\mathbf{v}$. So we can just rank by order of the v_j . This is the basis of an algorithm called the HITS algorithm, which was one of the early proposals for ranking web pages.

A different ranking called *page rank* is widely used. It is based on a random walk on the graph described above. We will study random walks in detail in Chapter 5.

4.5.5 An Application of SVD to a Discrete Optimization Problem

In the last example, the SVD was used as a dimension reduction technique. It found a k -dimensional subspace (the space of centers) of a d -dimensional space and made the Gaussian clustering problem easier by projecting the data to the subspace. Here, instead of fitting a model to data, we have an optimization problem. Again applying dimension reduction to the data makes the problem easier. The use of SVD to solve discrete optimization problems is a relatively new subject with many applications. We start with an important NP-hard problem, the maximum cut problem for a directed graph $G(V, E)$.

The maximum cut problem is to partition the node set V of a directed graph into two subsets S and $\bar{S}$ so that the number of edges from S to $\bar{S}$ is maximized. Let A be the adjacency matrix of the graph. With each vertex i , associate an indicator variable x_i . The variable x_i will be set to 1 for $i \in S$ and 0 for $i \in \bar{S}$. The vector $\mathbf{x} = (x_1, x_2, \dots, x_n)$ is unknown and we are trying to find it or equivalently the cut, so as to maximize the number of edges across the cut. The number of edges across the cut is precisely

$$\sum_{i,j} x_i (1 - x_j) a_{ij}.$$

Thus, the maximum cut problem can be posed as the optimization problem

$$\text{Maximize } \sum_{i,j} x_i (1 - x_j) a_{ij} \quad \text{subject to } x_i \in \{0, 1\}.$$

In matrix notation,

$$\sum_{i,j} x_i(1 - x_j)a_{ij} = \mathbf{x}^T A(\mathbf{1} - \mathbf{x}),$$

where $\mathbf{1}$ denotes the vector of all 1's. So, the problem can be restated as

$$\text{Maximize } \mathbf{x}^T A(\mathbf{1} - \mathbf{x}) \quad \text{subject to } x_i \in \{0, 1\}. \quad (4.1)$$

The SVD is used to solve this problem approximately by computing the SVD of A and replacing A by $A_k = \sum_{i=1}^k \sigma_i \mathbf{u}_i \mathbf{v}_i^T$ in (4.1) to get

$$\text{Maximize } \mathbf{x}^T A_k(\mathbf{1} - \mathbf{x}) \quad \text{subject to } x_i \in \{0, 1\}. \quad (4.2)$$

Note that the matrix A_k is no longer a 0-1 adjacency matrix.

We will show that:

1. For each 0-1 vector $\mathbf{x}$, $\mathbf{x}^T A_k(\mathbf{1} - \mathbf{x})$ and $\mathbf{x}^T A(\mathbf{1} - \mathbf{x})$ differ by at most $\frac{n^2}{\sqrt{k+1}}$. Thus, the maxima in (4.1) and (4.2) differ by at most this amount.
2. A near optimal $\mathbf{x}$ for (4.2) can be found by exploiting the low rank of A_k , which by Item 1 is near optimal for (4.1) where near optimal means with additive error of at most $\frac{n^2}{\sqrt{k+1}}$.

First, we prove Item 1. Since $\mathbf{x}$ and $\mathbf{1} - \mathbf{x}$ are 0-1 n -vectors, each has length at most $\sqrt{n}$. By the definition of the 2-norm, $|\mathbf{x}^T (A - A_k)(\mathbf{1} - \mathbf{x})| \leq \sqrt{n} \|A - A_k\|_2$. Now since $\mathbf{x}^T (A - A_k)(\mathbf{1} - \mathbf{x})$ is the dot product of the vector $\mathbf{x}$ with the vector $(A - A_k)(\mathbf{1} - \mathbf{x})$,

$$|\mathbf{x}^T (A - A_k)(\mathbf{1} - \mathbf{x})| \leq n \|A - A_k\|_2.$$

By Lemma 4.8, $\|A - A_k\|_2 = \sigma_{k+1}(A)$. The inequalities,

$$(k+1)\sigma_{k+1}^2 \leq \sigma_1^2 + \sigma_2^2 + \cdots + \sigma_{k+1}^2 \leq \|A\|_F^2 = \sum_{i,j} a_{ij}^2 \leq n^2$$

imply that $\sigma_{k+1}^2 \leq \frac{n^2}{k+1}$ and hence $\|A - A_k\|_2 \leq \frac{n}{\sqrt{k+1}}$ proving Item 1.

Next we focus on Item 2. It is instructive to look at the special case when $k=1$ and A is approximated by the rank one matrix A_1 . An even more special case when the left and right-singular vectors $\mathbf{u}$ and $\mathbf{v}$ are required to be identical is already NP-hard to solve exactly because it subsumes the problem of whether for a set of n integers, $\{a_1, a_2, \dots, a_n\}$, there is a partition into two subsets whose sums are equal. So, we look for algorithms that solve the maximum cut problem approximately.

For Item 2, we want to maximize $\sum_{i=1}^k \sigma_i (\mathbf{x}^T \mathbf{u}_i)(\mathbf{v}_i^T (\mathbf{1} - \mathbf{x}))$ over 0-1 vectors $\mathbf{x}$. A piece of notation will be useful. For any $S \subseteq \{1, 2, \dots, n\}$, write $\mathbf{u}_i(S)$ for the sum of

coordinates of the vector $\mathbf{u}_i$ corresponding to elements in the set S and also for $\mathbf{v}_i$. That is, $\mathbf{u}_i(S) = \sum_{j \in S} u_{ij}$. We will maximize $\sum_{i=1}^k \sigma_i \mathbf{u}_i(S) \mathbf{v}_i(\bar{S})$ using dynamic programming.

For a subset S of $\{1, 2, \dots, n\}$, define the $2k$ -dimensional vector

$$\mathbf{w}(S) = (\mathbf{u}_1(S), \mathbf{v}_1(\bar{S}), \mathbf{u}_2(S), \mathbf{v}_2(\bar{S}), \dots, \mathbf{u}_k(S), \mathbf{v}_k(\bar{S})).$$

If we had the list of all such vectors, we could find $\sum_{i=1}^k \sigma_i \mathbf{u}_i(S) \mathbf{v}_i(\bar{S})$ for each of them and take the maximum. There are 2^n subsets S , but several S could have the same $\mathbf{w}(S)$ and in that case it suffices to list just one of them. Round each coordinate of each $\mathbf{u}_i$ to the nearest integer multiple of $\frac{1}{nk^2}$. Call the rounded vector $\tilde{\mathbf{u}}_i$. Similarly obtain $\tilde{\mathbf{v}}_i$. Let $\tilde{\mathbf{w}}(S)$ denote the vector $(\tilde{\mathbf{u}}_1(S), \tilde{\mathbf{v}}_1(\bar{S}), \tilde{\mathbf{u}}_2(S), \tilde{\mathbf{v}}_2(\bar{S}), \dots, \tilde{\mathbf{u}}_k(S), \tilde{\mathbf{v}}_k(\bar{S}))$. We will construct a list of all possible values of the vector $\tilde{\mathbf{w}}(S)$. Again, if several different S 's lead to the same vector $\tilde{\mathbf{w}}(S)$, we will keep only one copy on the list. The list will be constructed by dynamic programming. For the recursive step of dynamic programming, assume we already have a list of all such vectors for $S \subseteq \{1, 2, \dots, i\}$ and wish to construct the list for $S \subseteq \{1, 2, \dots, i+1\}$. Each $S \subseteq \{1, 2, \dots, i\}$ leads to two possible $S' \subseteq \{1, 2, \dots, i+1\}$, namely, S and $S \cup \{i+1\}$. In the first case, the vector $\tilde{\mathbf{w}}(S') = (\tilde{\mathbf{u}}_1(S), \tilde{\mathbf{v}}_1(\bar{S}) + \tilde{v}_{1,i+1}, \tilde{\mathbf{u}}_2(S), \tilde{\mathbf{v}}_2(\bar{S}) + \tilde{v}_{2,i+1}, \dots, \dots)$. In the second case, it is $\tilde{\mathbf{w}}(S') = (\tilde{\mathbf{u}}_1(S) + \tilde{u}_{1,i+1}, \tilde{\mathbf{v}}_1(\bar{S}), \tilde{\mathbf{u}}_2(S) + \tilde{u}_{2,i+1}, \tilde{\mathbf{v}}_2(\bar{S}), \dots, \dots)$. We put in these two vectors for each vector in the previous list. Then, crucially, we prune - i.e., eliminate duplicates.

Assume that k is constant. Now, we show that the error is at most $\frac{n^2}{\sqrt{k+1}}$ as claimed. Since $\mathbf{u}_i, \mathbf{v}_i$ are unit length vectors, $|\mathbf{u}_i(S)|, |\mathbf{v}_i(\bar{S})| \leq \sqrt{n}$. Also $|\tilde{\mathbf{u}}_i(S) - \mathbf{u}_i(S)| \leq \frac{n}{nk^2} = \frac{1}{k^2}$ and similarly for $\mathbf{v}_i$. To bound the error, we use an elementary fact: if a, b are reals with $|a|, |b| \leq M$ and we estimate a by a' and b by b' so that $|a - a'|, |b - b'| \leq \delta \leq M$, then $a'b'$ is an estimate of ab in the sense

$$|ab - a'b'| = |a(b - b') + b'(a - a')| \leq |a||b - b'| + (|b| + |b - b'|)|a - a'| \leq 3M\delta.$$

Using this, we get that

$$\left| \sum_{i=1}^k \sigma_i \tilde{\mathbf{u}}_i(S) \tilde{\mathbf{v}}_i(\bar{S}) - \sum_{i=1}^k \sigma_i \mathbf{u}_i(S) \mathbf{v}_i(\bar{S}) \right| \leq 3k\sigma_1\sqrt{n}/k^2 \leq 3n^{3/2}/k \leq n^2/k,$$

and this meets the claimed error bound.

Next, we show that the running time is polynomially bounded. $|\tilde{\mathbf{u}}_i(S)|, |\tilde{\mathbf{v}}_i(S)| \leq 2\sqrt{n}$. Since $\tilde{\mathbf{u}}_i(S), \tilde{\mathbf{v}}_i(S)$ are all integer multiples of $1/(nk^2)$, there are at most $2/\sqrt{nk^2}$ possible values of $\tilde{\mathbf{u}}_i(S), \tilde{\mathbf{v}}_i(S)$ from which it follows that the list of $\tilde{\mathbf{w}}(S)$ never gets larger than $(1/\sqrt{nk^2})^{2k}$ which for fixed k is polynomially bounded.

We summarize what we have accomplished.

Theorem 4.17 *Given a directed graph $G(V, E)$, a cut of size at least the maximum cut minus $O\left(\frac{n^2}{\sqrt{k}}\right)$ can be computed in polynomial time n for any fixed k .*

It would be quite a surprise to have an algorithm that actually achieves the same accuracy in time polynomial in n and k because this would give an exact max cut in polynomial time.

4.6 Singular Vectors and Eigenvectors

An eigenvector of a square matrix A is a vector $\mathbf{v}$ satisfying $A\mathbf{v} = \lambda\mathbf{v}$, for a non-zero scalar λ which is the corresponding eigenvalue. A square matrix A can be viewed as a linear transformation from a space into itself which transforms an eigenvector into a scalar multiple of itself. The eigenvector decomposition of A is $V^T D V$ where the columns of V are the eigenvectors of A and D is a diagonal matrix with the eigenvalues on the diagonal.

A non square $m \times n$ matrix A also defines a linear transformation, but now from $\mathbf{R}^n$ to $\mathbf{R}^m$. In this case, eigenvectors do not make sense. But singular vectors can be defined. They serve the purpose of decomposing the linear transformation defined by the matrix A into the sum of simple linear transformations, each of which maps $\mathbf{R}_n$ to a one dimensional space, i.e., to a line through the origin.

A positive semi-definite matrix can be decomposed into a product AA^T . Thus, the eigenvector decomposition can be obtained from the singular value decomposition of $A = UDV^T$ since

$$AA^T = UDV^T VDU^T = UD^2U^T = \sum_i \sigma_i(A)^2 \mathbf{u}_i \mathbf{u}_i^T,$$

where the $\mathbf{u}_i$, the columns of U , are the eigenvectors of AA^T .

There are many applications of singular vectors and eigenvectors. For square non-symmetric matrices, both singular vectors and eigenvectors are defined but they may be different. In an important application, the pagerank, one represents the web by a $n \times n$ matrix A , where, a_{ij} is one if there is a hypertext link from the i^{th} page in the web to the j^{th} page. Otherwise, it is zero. The matrix is scaled by dividing each entry by the sum of entries in its row to get a stochastic matrix P . A stochastic matrix is one with nonnegative entries where each row sums to one. Note that P is not necessarily symmetric. Since the row sums of P are all one, the vector $\mathbf{1}$ of all one's is a right eigenvector with eigenvalue one, i.e., $P\mathbf{1} = \mathbf{1}$. This eigenvector contains no information. But the left eigenvector $\mathbf{v}$ with eigenvalue one satisfies $\mathbf{v}^T P = \mathbf{v}^T$ and is the stationary probability of the Markov chain with transition probability matrix P . So, it is the proportion of time a Markov chain spends at each vertex (page) in the long run. A simplified definition of pagerank ranks the page in order of its component in the top left eigenvector $\mathbf{v}$.

4.7 Bibliographic Notes

Singular value decomposition is fundamental to numerical analysis and linear algebra. There are many texts on these subjects and the interested reader may want to study these. A good reference is [GvL96]. The material on clustering a mixture of Gaussians in Section 4.5.2 is from [VW02]. Modeling data with a mixture of Gaussians is a standard tool in statistics. Several well-known heuristics like the expectation-minimization algorithm are used to learn (fit) the mixture model to data. Recently, in theoretical computer science, there has been modest progress on provable polynomial-time algorithms for learning mixtures. Some references are [DS07], [AK], [AM05], [MV10]. The application to the discrete optimization problem is from [FK99]. The section on ranking documents/webpages is from two influential papers, one on hubs and authorities by Jon Kleinberg [Kle99] and the other on pagerank by Page, Brin, Motwani and Winograd [BMPW98].

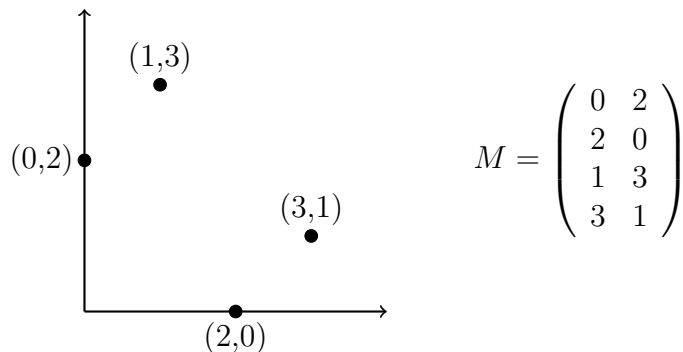


Figure 4.6: SVD problem

4.8 Exercises

Exercise 4.1 (Best fit functions versus best least squares fit) *In many experiments one collects the value of a parameter at various instances of time. Let y_i be the value of the parameter y at time x_i . Suppose we wish to construct the best linear approximation to the data in the sense that we wish to minimize the mean square error. Here error is measured vertically rather than perpendicular to the line. Develop formulas for m and b to minimize the mean square error of the points $\{(x_i, y_i) \mid 1 \leq i \leq n\}$ to the line $y = mx + b$.*

Exercise 4.2 *Given five observed parameters, height, weight, age, income, and blood pressure of n people, how would one find the best least squares fit subspace of the form*

$$a_1(\text{height}) + a_2(\text{weight}) + a_3(\text{age}) + a_4(\text{income}) + a_5(\text{blood pressure}) = 0$$

Here $a_1, a_2, \dots, a_5$ are the unknown parameters. If there is a good best fit 4-dimensional subspace, then one can think of the points as lying close to a 4-dimensional sheet rather than points lying in 5-dimensions. Why is it better to use the perpendicular distance to the subspace rather than vertical distance where vertical distance to the subspace is measured along the coordinate axis corresponding to one of the unknowns?

Exercise 4.3 *What is the best fit line through the origin for each of the following set of points?*

1. $\{(0, 1), (1, 0)\}$
2. $\{(0, 1), (2, 0)\}$
3. *The rows of the matrix*

$$\begin{pmatrix} 17 & 4 \\ -2 & 26 \\ 11 & 7 \end{pmatrix}$$

Exercise 4.4 Find the left and right-singular vectors, the singular values, and the SVD decomposition of the matrix M in Figure 4.6.

Exercise 4.5 Let

$$M = \begin{pmatrix} 1 & 1 \\ 0 & 3 \\ 3 & 0 \end{pmatrix}$$

Compute the first right-singular vector, the first singular value, and the first left-singular vector of M . Hint: first plot the rows of M on the 2-dimensional plane.

Exercise 4.6 Let A be a square $n \times n$ matrix whose rows are orthonormal. Prove that the columns of A are orthonormal.

Exercise 4.7 Suppose A is a $n \times n$ matrix with block diagonal structure with k equal size blocks where all entries of the i^{th} block are a_i with $a_1 > a_2 > \dots > a_k > 0$. Show that A has exactly k nonzero singular vectors $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_k$ where $\mathbf{v}_i$ has the value $(\frac{k}{n})^{1/2}$ in the coordinates corresponding to the i^{th} block and 0 elsewhere. In other words, the singular vectors exactly identify the blocks of the diagonal. What happens if $a_1 = a_2 = \dots = a_k$? In the case where the a_i are equal, what is the structure of the set of all possible singular vectors?

Hint: By symmetry, the top singular vector's components must be constant in each block.

Exercise 4.8 Prove that the left-singular vectors of A are the right-singular vectors of A^T .

Exercise 4.9 Interpret the first right and left-singular vectors for the document term matrix.

Exercise 4.10 Verify that the sum of r -rank one matrices $\sum_{i=1}^r \sigma_i \mathbf{u}_i \mathbf{v}_i^T$ can be written as UDV^T , where the $\mathbf{u}_i$ are the columns of U and $\mathbf{v}_i$ are the columns of V . To do this, first verify that for any two matrices P and Q , we have

$$PQ^T = \sum_i \mathbf{p}_i \mathbf{q}_i^T$$

where $\mathbf{p}_i$ is the i^{th} column of P and $\mathbf{q}_i$ is the i^{th} column of Q .

Exercise 4.11

1. Show that the rank of A is r where r is the minimum i such that $\arg \max_{\substack{\mathbf{v} \perp \mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_i \\ |\mathbf{v}|=1}} |A \mathbf{v}| = 0$.

2. Show that $|\mathbf{u}_1^T A| = \max_{|\mathbf{u}|=1} |\mathbf{u}^T A| = \sigma_1$.

Hint: Use SVD.

Exercise 4.12 If $\sigma_1, \sigma_2, \dots, \sigma_r$ are the singular values of A and $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r$ are the corresponding right-singular vectors, show that

1. $A^T A = \sum_{i=1}^r \sigma_i^2 \mathbf{v}_i \mathbf{v}_i^T$
2. $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r$ are eigenvectors of $A^T A$.
3. Assuming that the set of eigenvectors of a matrix is unique, conclude that the set of singular values of the matrix is unique.

See the appendix for the definition of eigenvectors.

Exercise 4.13 Let $\sum_i \sigma_i u_i v_i^T$ be the singular value decomposition of a rank r matrix A .

Let $A_k = \sum_{i=1}^k \sigma_i u_i v_i^T$ be a rank k approximation to A . Express the following quantities in terms of the singular values $\{\sigma_i, 1 \leq i \leq r\}$.

1. $\|A_k\|_F^2$
2. $\|A_k\|_2^2$
3. $\|A - A_k\|_F^2$
4. $\|A - A_k\|_2^2$

Exercise 4.14 If A is a symmetric matrix with distinct singular values, show that the left and right singular vectors are the same and that $A = V D V^T$.

Exercise 4.15 Let A be a matrix. Given an algorithm for finding

$$\mathbf{v}_1 = \arg \max_{|\mathbf{v}|=1} |A\mathbf{v}|,$$

describe an algorithm to find the SVD of A .

Exercise 4.16 Compute the singular valued decomposition of the matrix

$$A = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix}$$

Exercise 4.17 Write a program to implement the power method for computing the first singular vector of a matrix. Apply your program to the matrix

$$A = \begin{pmatrix} 1 & 2 & 3 & \cdots & 9 & 10 \\ 2 & 3 & 4 & \cdots & 10 & 0 \\ \vdots & \vdots & \vdots & & & \vdots \\ 9 & 10 & 0 & \cdots & 0 & 0 \\ 10 & 0 & 0 & \cdots & 0 & 0 \end{pmatrix}.$$

Exercise 4.18 Modify the power method to find the first four singular vectors of a matrix A as follows. Randomly select four vectors and find an orthonormal basis for the space spanned by the four vectors. Then multiply each of the basis vectors times A and find a new orthonormal basis for the space spanned by the resulting four vectors. Apply your method to find the first four singular vectors of matrix A of Exercise 4.17

Exercise 4.19 Let A be a real valued matrix. Prove that $B = AA^T$ is positive semi definite. A matrix B is positive semi definite if for all $\mathbf{x}$, $\mathbf{x}^T B \mathbf{x} \geq 0$.

Exercise 4.20 Let A be the adjacency matrix of a graph. The Laplacian of A is $L = D - A$ where D is a diagonal matrix whose diagonal entries are the row sums of A . Prove that L is positive semi definite. A matrix L is positive semi definite if for all $\mathbf{x}$, $\mathbf{x}^T L \mathbf{x} \geq 0$.

Exercise 4.21 Prove that the eigenvalues of a symmetric real valued matrix are real.

Exercise 4.22 Suppose A is a square invertible matrix and the SVD of A is $A = \sum_i \sigma_i u_i v_i^T$. Prove that the inverse of A is $\sum_i \frac{1}{\sigma_i} v_i u_i^T$.

Exercise 4.23 Suppose A is square, but not necessarily invertible and has SVD $A = \sum_{i=1}^r \sigma_i u_i v_i^T$. Let $B = \sum_{i=1}^r \frac{1}{\sigma_i} v_i u_i^T$. Show that $BA\mathbf{x} = \mathbf{x}$ for all $\mathbf{x}$ in the span of the right-singular vectors of A . For this reason B is sometimes called the pseudo inverse of A and can play the role of A^{-1} in many applications.

Exercise 4.24

1. For any matrix A , show that $\sigma_k \leq \frac{\|A\|_F}{\sqrt{k}}$.
2. Prove that there exists a matrix B of rank at most k such that $\|A - B\|_2 \leq \frac{\|A\|_F}{\sqrt{k}}$.
3. Can the 2-norm on the left hand side in (b) be replaced by Frobenius norm?

Exercise 4.25 Suppose an $n \times d$ matrix A is given and you are allowed to preprocess A . Then you are given a number of d -dimensional vectors $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m$ and for each of these vectors you must find the vector $A\mathbf{x}_i$ approximately, in the sense that you must find a vector $\mathbf{u}_i$ satisfying $|\mathbf{u}_i - A\mathbf{x}_i| \leq \varepsilon \|A\|_F |\mathbf{x}_i|$. Here $\varepsilon > 0$ is a given error bound. Describe an algorithm that accomplishes this in time $O\left(\frac{d+n}{\varepsilon^2}\right)$ per $\mathbf{x}_i$ not counting the preprocessing time.

Exercise 4.26 (Constrained Least Squares Problem using SVD) Given A , $\mathbf{b}$, and m , use the SVD algorithm to find a vector $\mathbf{x}$ with $|\mathbf{x}| < m$ minimizing $|\mathbf{Ax} - \mathbf{b}|$. This problem is a learning exercise for the advanced student. For hints/solution consult Golub and van Loan, Chapter 12.

Exercise 4.27 (Document-Term Matrices): Suppose we have a $m \times n$ document-term matrix where each row corresponds to a document where the rows have been normalized to length one. Define the “similarity” between two such documents by their dot product.

1. Consider a “synthetic” document whose sum of squared similarities with all documents in the matrix is as high as possible. What is this synthetic document and how would you find it?
2. How does the synthetic document in (1) differ from the center of gravity?
3. Building on (1), given a positive integer k , find a set of k synthetic documents such that the sum of squares of the mk similarities between each document in the matrix and each synthetic document is maximized. To avoid the trivial solution of selecting k copies of the document in (1), require the k synthetic documents to be orthogonal to each other. Relate these synthetic documents to singular vectors.
4. Suppose that the documents can be partitioned into k subsets (often called clusters), where documents in the same cluster are similar and documents in different clusters are not very similar. Consider the computational problem of isolating the clusters. This is a hard problem in general. But assume that the terms can also be partitioned into k clusters so that for $i \neq j$, no term in the i^{th} cluster occurs in a document in the j^{th} cluster. If we knew the clusters and arranged the rows and columns in them to be contiguous, then the matrix would be a block-diagonal matrix. Of course the clusters are not known. By a “block” of the document-term matrix, we mean a submatrix with rows corresponding to the i^{th} cluster of documents and columns corresponding to the i^{th} cluster of terms. We can also partition any n vector into blocks. Show that any right-singular vector of the matrix must have the property that each of its blocks is a right-singular vector of the corresponding block of the document-term matrix.
5. Suppose now that the singular values of all the blocks are distinct (also across blocks). Show how to solve the clustering problem.

Hint: (4) Use the fact that the right-singular vectors must be eigenvectors of $A^T A$. Show that $A^T A$ is also block-diagonal and use properties of eigenvectors.

Exercise 4.28 Generate a number of samples according to a mixture of 1-dimensional Gaussians. See what happens as the centers get closer. Alternatively, see what happens when the centers are fixed and the standard deviation is increased.

Exercise 4.29 (Newcomb/Binford) The frequency distribution of first digits in many data sets is not uniform. One might expect the distribution to be scale free in that changing the units of measure should not change the distribution. Determine a distribution where multiplying each number by two does not change the distribution. Hint: Construct a graph with nine vertices where each vertex corresponds to one of the nine first digits. The edge

from vertex i to vertex j is labeled with the probability that multiplying a number whose first digit is i by 2 results in a number whose first digit is j . What is the stationary probability of a random walk on this graph?

For scale factor 3 the adjacency matrix is

$$\begin{pmatrix} 0 & 0 & \frac{1}{3} & \frac{1}{3} & \frac{1}{3} & 0 & 0 & 0 & 0 \\ \frac{2}{3} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & \frac{1}{3} \\ 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ \frac{2}{3} & \frac{1}{3} & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

The first eigenvector of A^T is

0.3214 0.1607 0.1071 0.1071 0.0536 0.0536 0.0536 0.0536 0.0357

Exercise 4.30 Show that maximizing $\mathbf{x}^T \mathbf{u} \mathbf{u}^T (\mathbf{1} - \mathbf{x})$ subject to $x_i \in \{0, 1\}$ is equivalent to partitioning the coordinates of $\mathbf{u}$ into two subsets where the sum of the elements in both subsets are equal.

Exercise 4.31 Read in a photo and convert to a matrix. Perform a singular value decomposition of the matrix. Reconstruct the photo using only 10%, 25%, 50% of the singular values.

1. Print the reconstructed photo. How good is the quality of the reconstructed photo?
2. What percent of the Frobenius norm is captured in each case?

Hint: If you use Matlab, the command to read a photo is `imread`. The types of files that can be read are given by `imformats`. To print the file use `imwrite`. Print using jpeg format. To access the file afterwards you may need to add the file extension `.jpg`. The command `imread` will read the file in `uint8` and you will need to convert to `double` for the SVD code. Afterwards you will need to convert back to `uint8` to write the file. If the photo is a color photo you will get three matrices for the three colors used.

Exercise 4.32 Create a set of 100, 100×100 matrices of random numbers between 0 and 1 such that each entry is highly correlated with the adjacency entries. Find the SVD of A . What fraction of the Frobenius norm of A is captured by the top 10 singular vectors? How many singular vectors are required to capture 95% of the Frobenius norm?

Exercise 4.33 Create a 100×100 matrix A of random numbers between 0 and 1 such that each entry is highly correlated with the adjacency entries and find the first 10 vectors for a single basis that is reasonably good for all 100 matrices. How does one do this? What fraction of the Frobenius norm of a new matrix is captured by the basis?

Exercise 4.34 Show that the running time for the maximum cut algorithm in Section ?? can be carried out in time $O(n^3 + \text{poly}(n)k^k)$, where poly is some polynomial.

Exercise 4.35 Let $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ be n points in d -dimensional space and let X be the $n \times d$ matrix whose rows are the n points. Suppose we know only the matrix D of pairwise distances between points and not the coordinates of the points themselves. The set of points $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ giving rise to the distance matrix D is not unique since any translation, rotation, or reflection of the coordinate system leaves the distances invariant. Fix the origin of the coordinate system so that the centroid of the set of points is at the origin. That is, $\sum_{i=1}^n \mathbf{x}_i = 0$.

1. Show that the elements of XX^T are given by

$$\mathbf{x}_i \mathbf{x}_j^T = -\frac{1}{2} \left[d_{ij}^2 - \frac{1}{n} \sum_{j=1}^n d_{ij}^2 - \frac{1}{n} \sum_{i=1}^n d_{ij}^2 + \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n d_{ij}^2 \right].$$

2. Describe an algorithm for determining the matrix X whose rows are the $\mathbf{x}_i$.

Exercise 4.36

1. Consider the pairwise distance matrix for twenty US cities given below. Use the algorithm of Exercise 4.35 to place the cities on a map of the US. The algorithm is called classical multidimensional scaling, `cmdscale`, in Matlab. Alternatively use the pairwise distance matrix to place the cities on a map of China.

Note: Any rotation or a mirror image of the map will have the same pairwise distances.

2. Suppose you had airline distances for 50 cities around the world. Could you use these distances to construct a world map?

	B	B	C	D	D	H	L	M	M	M
	O	U	H	A	E	O	A	E	I	I
	S	F	I	L	N	U		M	A	M
Boston	-	400	851	1551	1769	1605	2596	1137	1255	1123
Buffalo	400	-	454	1198	1370	1286	2198	803	1181	731
Chicago	851	454	-	803	920	940	1745	482	1188	355
Dallas	1551	1198	803	-	663	225	1240	420	1111	862
Denver	1769	1370	920	663	-	879	831	879	1726	700
Houston	1605	1286	940	225	879	-	1374	484	968	1056
Los Angeles	2596	2198	1745	1240	831	1374	-	1603	2339	1524
Memphis	1137	803	482	420	879	484	1603	-	872	699
Miami	1255	1181	1188	1111	1726	968	2339	872	-	1511
Minneapolis	1123	731	355	862	700	1056	1524	699	1511	-
New York	188	292	713	1374	1631	1420	2451	957	1092	1018
Omaha	1282	883	432	586	488	794	1315	529	1397	290
Philadelphia	271	279	666	1299	1579	1341	2394	881	1019	985
Phoenix	2300	1906	1453	887	586	1017	357	1263	1982	1280
Pittsburgh	483	178	410	1070	1320	1137	2136	660	1010	743
Saint Louis	1038	662	262	547	796	679	1589	240	1061	466
Salt Lake City	2099	1699	1260	999	371	1200	579	1250	2089	987
San Francisco	2699	2300	1858	1483	949	1645	347	1802	2594	1584
Seattle	2493	2117	1737	1681	1021	1891	959	1867	2734	1395
Washington D.C.	393	292	597	1185	1494	1220	2300	765	923	934

	N Y	O M A	P H I	P H O	P I T	S t L	S L C	S F	S E A	D C
Boston	188	1282	271	2300	483	1038	2099	2699	2493	393
Buffalo	292	883	279	1906	178	662	1699	2300	2117	292
Chicago	713	432	666	1453	410	262	1260	1858	1737	597
Dallas	1374	586	1299	887	1070	547	999	1483	1681	1185
Denver	1631	488	1579	586	1320	796	371	949	1021	1494
Houston	1420	794	1341	1017	1137	679	1200	1645	1891	1220
Los Angeles	2451	1315	2394	357	2136	1589	579	347	959	2300
Memphis	957	529	881	1263	660	240	1250	1802	1867	765
Miami	1092	1397	1019	1982	1010	1061	2089	2594	2734	923
Minneapolis	1018	290	985	1280	743	466	987	1584	1395	934
New York	-	1144	83	2145	317	875	1972	2571	2408	230
Omaha	1144	-	1094	1036	836	354	833	1429	1369	1014
Philadelphia	83	1094	-	2083	259	811	1925	2523	2380	123
Phoenix	2145	1036	2083	-	1828	1272	504	653	1114	1973
Pittsburgh	317	836	259	1828	-	559	1668	2264	2138	192
Saint Louis	875	354	811	1272	559	-	1162	1744	1724	712
Salt Lake City	1972	833	1925	504	1668	1162	-	600	701	1848
San Francisco	2571	1429	2523	653	2264	1744	600	-	678	2442
Seattle	2408	1369	2380	1114	2138	1724	701	678	-	2329
Washington D.C.	230	1014	123	1973	192	712	1848	2442	2329	-

City	Bei- jing	Tian- jin	Shang- hai	Chong- qing	Hoh- hot	Urum- qi	Lha- sa	Yin- chuan	Nan- ning	Har- bin	Chang- chun	Shen- yang
Beijing	0	125	1239	3026	480	3300	3736	1192	2373	1230	979	684
Tianjin	125	0	1150	1954	604	3330	3740	1316	2389	1207	955	661
Shanghai	1239	1150	0	1945	1717	3929	4157	2092	1892	2342	2090	1796
Chongqing	3026	1954	1945	0	1847	3202	2457	1570	993	3156	2905	2610
Hohhot	480	604	1717	1847	0	2825	3260	716	2657	1710	1458	1164
Urumqi	3300	3330	3929	3202	2825	0	2668	2111	4279	4531	4279	3985
Lhasa	3736	3740	4157	2457	3260	2668	0	2547	3431	4967	4715	4421
Yinchuan	1192	1316	2092	1570	716	2111	2547	0	2673	2422	2170	1876
Nanning	2373	2389	1892	993	2657	4279	3431	2673	0	3592	3340	3046
Harbin	1230	1207	2342	3156	1710	4531	4967	2422	3592	0	256	546
Changchun	979	955	2090	2905	1458	4279	4715	2170	3340	256	0	294
Shenyang	684	661	1796	2610	1164	3985	4421	1876	3046	546	294	0

5 Random Walks and Markov Chains

A random walk on a directed graph consists of a sequence of vertices generated from a start vertex by selecting an edge, traversing the edge to a new vertex, and repeating the process. We will see that if the graph is strongly connected, then the fraction of time the walk spends at the various vertices of the graph converges to a stationary probability distribution.

Since the graph is directed, there might be vertices with no out edges and hence nowhere for the walk to go. Vertices in a strongly connected component with no in edges from the remainder of the graph can never be reached unless the component contains the start vertex. Once a walk leaves a strongly connected component it can never return. Most of our discussion of random walks will involve strongly connected graphs.

Start a random walk at a vertex x_0 and think of the starting probability distribution as putting a mass of one on x_0 and zero on every other vertex. More generally, one could start with any probability distribution $\mathbf{p}$, where $\mathbf{p}$ is a row vector with nonnegative components summing to one, with p_x being the probability of starting at vertex x . The probability of being at vertex x at time $t + 1$ is the sum over each adjacent vertex y of being at y at time t and taking the transition from y to x . Let $\mathbf{p}^{(t)}$ be a row vector with a component for each vertex specifying the probability mass of the vertex at time t and let $\mathbf{p}^{(t+1)}$ be the row vector of probabilities at time $t + 1$. In matrix notation³

$$\mathbf{p}^{(t)} P = \mathbf{p}^{(t+1)}$$

where the ij^{th} entry of the matrix P is the probability of the walk at vertex i selecting the edge to vertex j .

A fundamental property of a random walk is that in the limit, the long-term average probability of being at a particular vertex is independent of the start vertex, or an initial probability distribution over vertices, provided only that the underlying graph is strongly connected. The limiting probabilities are called the *stationary probabilities*. This fundamental theorem is proved in the next section.

A special case of random walks, namely random walks on undirected graphs, has important connections to electrical networks. Here, each edge has a parameter called *conductance*, like the electrical conductance, and if the walk is at vertex u , it chooses the edge from among all edges incident to u to walk to the next vertex with probabilities proportional to their conductance. Certain basic quantities associated with random walks are hitting time, the expected time to reach vertex y starting at vertex x , and cover time, the expected time to visit every vertex. Qualitatively, these quantities are all bounded above by polynomials in the number of vertices. The proofs of these facts will rely on the

³Probability vectors are represented by row vectors to simplify notation in equations like the one here.

random walk	Markov chain
graph	stochastic process
vertex	state
strongly connected	persistent
aperiodic	aperiodic
strongly connected and aperiodic	ergodic
undirected graph	time reversible

Table 5.1: Correspondence between terminology of random walks and Markov chains

analogy between random walks and electrical networks.

Aspects of the theory of random walks was developed in computer science with an important application in defining the pagerank of pages on the World Wide Web by their stationary probability. An equivalent concept called a *Markov chain* had previously been developed in the statistical literature. A Markov chain has a finite set of *states*. For each pair of states x and y , there is a *transition probability* p_{xy} of going from state x to state y where for each x , $\sum_y p_{xy} = 1$. A random walk in the Markov chain starts at some state. At a given time step, if it is in state x , the next state y is selected randomly with probability p_{xy} . A Markov chain can be represented by a directed graph with a vertex representing each state and an edge with weight p_{xy} from vertex x to vertex y . We say that the Markov chain is *connected* if the underlying directed graph is strongly connected. That is, if there is a directed path from every vertex to every other vertex. The matrix P consisting of the p_{xy} is called the *transition probability matrix* of the chain. The terms “random walk” and “Markov chain” are used interchangeably. The correspondence between the terminologies of random walks and Markov chains is given in Table 5.1.

A state of a Markov chain is *persistent* if it has the property that should the state ever be reached, the random process will return to it with probability one. This is equivalent to the property that the state is in a strongly connected component with no out edges. For most of the chapter, we assume that the underlying directed graph is strongly connected. We discuss here briefly what might happen if we do not have strong connectivity. Consider the directed graph in Figure 5.1b with three strongly connected components, A , B , and C . Starting from any vertex in A , there is a nonzero probability of eventually reaching any vertex in A . However, the probability of returning to a vertex in A is less than one and thus vertices in A , and similarly vertices in B , are not persistent. From any vertex in C , the walk eventually will return with probability one to the vertex, since there is no way of leaving component C . Thus, vertices in C are persistent.

Markov chains are used to model situations where all the information of the system

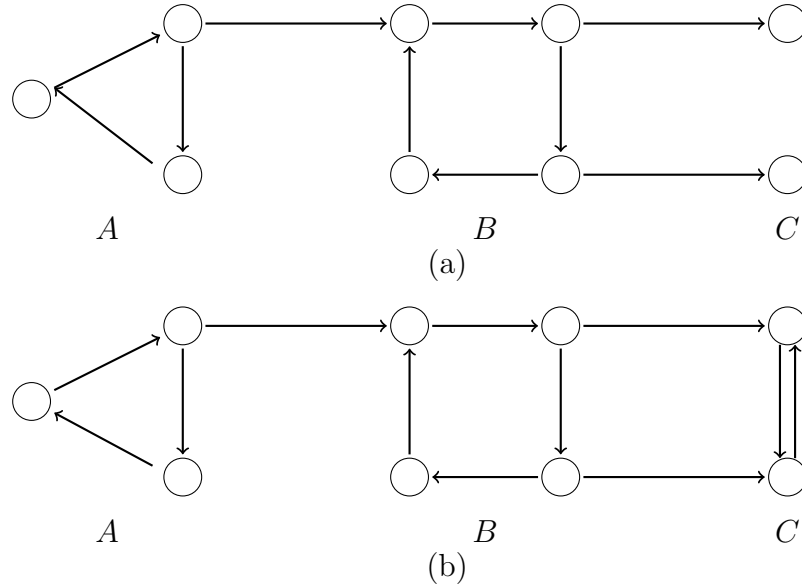


Figure 5.1: (a) A directed graph with vertices having no out edges and a strongly connected component *A* with no in edges.
(b) A directed graph with three strongly connected components.

necessary to predict the future can be encoded in the current state. A typical example is speech, where for a small k the current state encodes the last k syllables uttered by the speaker. Given the current state, there is a certain probability of each syllable being uttered next and these can be used to calculate the transition probabilities. Another example is a gambler's assets, which can be modeled as a Markov chain where the current state is the amount of money the gambler has on hand. The model would only be valid if the gambler's bets depend only on current assets, not the past.

Later in the chapter, we study the widely used Markov Chain Monte Carlo method (MCMC). Here, the objective is to sample a large space according to some probability distribution p . The number of elements in the space may be very large, say 10^{100} . One designs a Markov chain where states correspond to the elements of the space. The transition probabilities of the chain are designed so that the stationary probability of the chain is the probability distribution p with which we want to sample. One samples by taking a random walk until the probability distribution is close to the stationary distribution of the chain and then selects the point the walk is at. The walk continues a number of steps until the probability distribution is no longer dependent on where the walk was when the first element was selected. A second point is then selected, and so on. Although it is impossible to store the graph in a computer since it has 10^{100} vertices, to do the walk one needs only store the vertex the walk is at and be able to generate the adjacent vertices by some algorithm. What is critical is that the probability of the walk converges to the stationary probability in time logarithmic in the number of states.

We mention two motivating examples. The first is to estimate the probability of a region R in d -space according to a probability density like the Gaussian. Put down a grid and make each grid point that is in R a state of the Markov chain. Given a probability density p , design transition probabilities of a Markov chain so that the stationary distribution is exactly p . In general, the number of states grows exponentially in the dimension d , but the time to converge to the stationary distribution grows polynomially in d .

A second example is from physics. Consider an $n \times n$ grid in the plane with a particle at each grid point. Each particle has a spin of ± 1 . There are 2^{n^2} spin configurations. The energy of a configuration is a function of the spins. A central problem in statistical mechanics is to sample a spin configuration according to their probability. It is easy to design a Markov chain with one state per spin configuration so that the stationary probability of a state is proportional to the state's energy. If a random walk gets close to the stationary probability in time polynomial to n rather than 2^{n^2} , then one can sample spin configurations according to their probability.

A quantity called the *mixing time*, loosely defined as the time needed to get close to the stationary distribution, is often much smaller than the number of states. In Section 5.8, we relate the mixing time to a combinatorial notion called *normalized conductance* and derive good upper bounds on the mixing time in many cases.

5.1 Stationary Distribution

Let $\mathbf{p}^{(t)}$ be the probability distribution after t steps of a random walk. Define the *long-term probability distribution* $\mathbf{a}^{(t)}$ by

$$\mathbf{a}^{(t)} = \frac{1}{t} (\mathbf{p}^{(0)} + \mathbf{p}^{(1)} + \cdots + \mathbf{p}^{(t-1)}).$$

The fundamental theorem of Markov chains asserts that the long-term probability distribution of a connected Markov chain converges to a unique limit probability vector, which we denote by $\boldsymbol{\pi}$. Executing one more step, starting from this limit distribution, we get back the same distribution. In matrix notation, $\boldsymbol{\pi}P = \boldsymbol{\pi}$ where P is the matrix of transition probabilities. In fact, there is a unique probability vector (nonnegative components summing to one) satisfying $\boldsymbol{\pi}P = \boldsymbol{\pi}$ and this vector is the limit. Also since one step does not change the distribution, any number of steps would not either. For this reason, $\boldsymbol{\pi}$ is called the *stationary distribution*.

Before proving the fundamental theorem of Markov chains, we first prove a technical lemma.

Lemma 5.1 *Let P be the transition probability matrix for a connected Markov chain. The $n \times (n+1)$ matrix $A = [P - I, \mathbf{1}]$ obtained by augmenting the matrix $P - I$ with an additional column of ones has rank n .*

Proof: If the rank of $A = [P - I, \mathbf{1}]$ was less than n there would be two linearly independent solutions to $A\mathbf{x} = \mathbf{0}$. Each row in P sums to one so each row in $P - I$ sums to zero. Thus $\mathbf{x} = (\mathbf{1}, 0)$, where all but the last coordinate of $\mathbf{x}$ is 1, is one solution to $A\mathbf{x} = \mathbf{0}$. Assume there was a second solution $(\mathbf{x}, \alpha)$ perpendicular to $(\mathbf{1}, 0)$. Then $(P - I)\mathbf{x} + \alpha\mathbf{1} = \mathbf{0}$ or $x_i = \sum_j p_{ij}x_j + \alpha$. Each x_i is a convex combination of some x_j plus α . Let S be the set of i for which x_i attains its maximum value. $\bar{S}$ is not empty since x is perpendicular to $\mathbf{1}$ and hence $\sum_j x_j = 0$. Connectedness implies that some x_k of maximum value is adjacent to some x_l of lower value. Thus, $x_k > \sum_j p_{kj}x_j$. Therefore α must be greater than 0 in $x_k = \sum_j p_{kj}x_j + \alpha$.

A symmetric argument with T the set of i with x_i taking its minimum value implies $\alpha < 0$ producing a contradiction thereby proving the lemma. ■

Theorem 5.2 (Fundamental Theorem of Markov Chains) *For a connected Markov chain there is a unique probability vector $\boldsymbol{\pi}$ satisfying $\boldsymbol{\pi}P = \boldsymbol{\pi}$. Moreover, for any starting distribution, $\lim_{t \rightarrow \infty} \mathbf{a}^{(t)}$ exists and equals $\boldsymbol{\pi}$.*

Proof: Note that $\mathbf{a}^{(t)}$ is itself a probability vector, since its components are nonnegative and sum to 1. Run one step of the Markov chain starting with distribution $\mathbf{a}^{(t)}$; the distribution after the step is $\mathbf{a}^{(t)}P$. Calculate the change in probabilities due to this step.

$$\begin{aligned} \mathbf{a}^{(t)}P - \mathbf{a}^{(t)} &= \frac{1}{t} [\mathbf{p}^{(0)}P + \mathbf{p}^{(1)}P + \cdots + \mathbf{p}^{(t-1)}P] - \frac{1}{t} [\mathbf{p}^{(0)} + \mathbf{p}^{(1)} + \cdots + \mathbf{p}^{(t-1)}] \\ &= \frac{1}{t} [\mathbf{p}^{(1)} + \mathbf{p}^{(2)} + \cdots + \mathbf{p}^{(t)}] - \frac{1}{t} [\mathbf{p}^{(0)} + \mathbf{p}^{(1)} + \cdots + \mathbf{p}^{(t-1)}] \\ &= \frac{1}{t} (\mathbf{p}^{(t)} - \mathbf{p}^{(0)}). \end{aligned}$$

Thus, $\mathbf{b}^{(t)} = \mathbf{a}^{(t)}P - \mathbf{a}^{(t)}$ satisfies $|\mathbf{b}^{(t)}| \leq \frac{2}{t} \rightarrow 0$, as $t \rightarrow \infty$.

By Lemma 5.1 above, A has rank n . The $n \times n$ submatrix B of A consisting of all its columns except the first is invertible. Let $\mathbf{c}^{(t)}$ be obtained from $\mathbf{b}^{(t)}$ by removing the first entry. Then, $\mathbf{a}^{(t)}B = [\mathbf{c}^{(t)}, 1]$ and so $\mathbf{a}^{(t)} = [\mathbf{c}^{(t)}, 1]B^{-1} \rightarrow [\mathbf{0}, 1]B^{-1}$. We have the theorem with $\boldsymbol{\pi} = [\mathbf{0}, 1]B^{-1}$. ■

Observe that the expected time r_x for a Markov chain starting in state x to return to state x is the reciprocal of the stationary probability of x . That is $r_x = \frac{1}{\pi_x}$. Intuitively this follows by observing that if a long walk always returns to state x in exactly r_x steps, the frequency of being in a state x would be $\frac{1}{r_x}$. A rigorous proof requires the Strong Law of Large Numbers.

We finish this section with the following lemma useful in establishing that a probability distribution is the stationary probability distribution for a random walk on a connected, directed graph.

Lemma 5.3 *For a random walk on a strongly connected, directed graph, if the vector π satisfies $\pi_x p_{xy} = \pi_y p_{yx}$ for all x and y and $\sum_x \pi_x = 1$, then π is the stationary distribution of the walk.*

Proof: Since π satisfies $\pi_x p_{xy} = \pi_y p_{yx}$, take the sum of both sides to get $\pi_x = \sum_y \pi_y p_{yx}$ and hence π satisfies $\pi = \pi P$. By Theorem 5.2, π is the unique stationary probability. ■

5.2 Electrical Networks and Random Walks

In the next few sections, we study the relationship between electrical networks and random walks on undirected graphs. The graphs have nonnegative weights on each edge. A step is executed by picking a random edge from the current vertex with probability proportional to the edge's weight and traversing the edge.

An electrical network is a connected, undirected graph in which each edge (x, y) has a resistance $r_{xy} > 0$. In what follows, it is easier to deal with conductance defined as the reciprocal of resistance, $c_{xy} = \frac{1}{r_{xy}}$, rather than resistance. Associated with an electrical network is a random walk on the underlying graph defined by assigning a probability $p_{xy} = \frac{c_{xy}}{c_x}$ to the edge (x, y) incident to the vertex x , where the normalizing constant c_x equals $\sum_y c_{xy}$. Note that although c_{xy} equals c_{yx} , the probabilities p_{xy} and p_{yx} may not be equal due to the normalization required to make the probabilities at each vertex sum to one. We shall soon see that there is a relationship between current flowing in an electrical network and a random walk on the underlying graph.

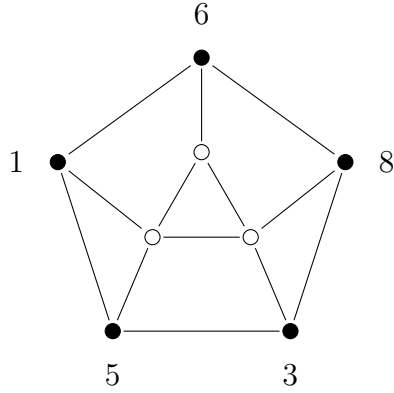
Since we assume that the undirected graph is connected, by Theorem 5.2 there is a unique stationary probability distribution. The stationary probability distribution is π where $\pi_x = \frac{c_x}{c_0}$ where $c_0 = \sum_x c_x$. To see this, for all x and y

$$\pi_x p_{xy} = \frac{c_x}{c_0} \frac{c_{xy}}{c_x} = \frac{c_{xy}}{c_0} = \frac{c_y}{c_0} \frac{c_{yx}}{c_y} = \pi_y p_{yx}$$

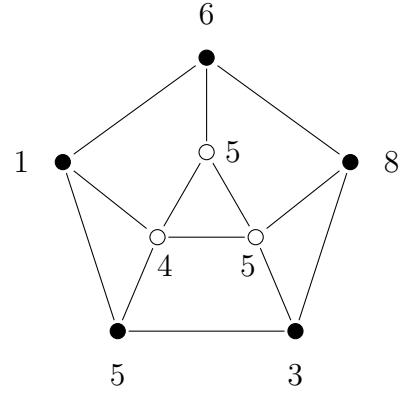
and hence by Lemma 5.3, π is the unique stationary probability.

Harmonic functions

Harmonic functions are useful in developing the relationship between electrical networks and random walks on undirected graphs. Given an undirected graph, designate a nonempty set of vertices as boundary vertices and the remaining vertices as interior vertices. A harmonic function g on the vertices is one in which the value of the function at the boundary vertices is fixed to some boundary condition and the value of g at any interior vertex x is a weighted average of the values at all the adjacent vertices y , with



Graph with boundary vertices dark and boundary conditions specified.



Values of harmonic function satisfying boundary conditions

Figure 5.2: Graph illustrating an harmonic function.

weights p_{xy} satisfying $\sum_y p_{xy} = 1$ for each x . Thus, if at every interior vertex x for some set of weights p_{xy} satisfying $\sum_y p_{xy} = 0$, $g_x = \sum_y g_y p_{xy}$, then g is an harmonic function.

Example: Convert an electrical network with conductances c_{xy} to a weighted, undirected graph with probabilities p_{xy} . Let $\mathbf{f}$ be a function satisfying $\mathbf{f}P = \mathbf{f}$ where P is the matrix of probabilities. It follows that the function $g_x = \frac{f_x}{c_x}$ is harmonic.

$$\begin{aligned} g_x &= \frac{f_x}{c_x} = \frac{1}{c_x} \sum_y f_y p_{yx} = \frac{1}{c_x} \sum_y f_y \frac{c_{yx}}{c_y} \\ &= \frac{1}{c_x} \sum_y f_y \frac{c_{xy}}{c_y} = \sum_y \frac{f_y}{c_y} \frac{c_{xy}}{c_x} = \sum_y g_y p_{xy} \end{aligned}$$

■

A harmonic function on a connected graph takes on its maximum and minimum on the boundary. Suppose the maximum does not occur on the boundary. Let S be the set of interior vertices at which the maximum value is attained. Since S contains no boundary vertices, $\bar{S}$ is nonempty. Connectedness implies that there is at least one edge (x, y) with $x \in S$ and $y \in \bar{S}$. The value of the function at x is the average of the value at its neighbors, all of which are less than or equal to the value at x and the value at y is strictly less, a contradiction. The proof for the minimum value is identical.

There is at most one harmonic function satisfying a given set of equations and boundary conditions. For suppose there were two solutions, $f(x)$ and $g(x)$. The difference of two solutions is itself harmonic. Since $h(x) = f(x) - g(x)$ is harmonic and has value zero on the boundary, by the min and max principles it has value zero everywhere. Thus $f(x) = g(x)$.

The analogy between electrical networks and random walks

There are important connections between electrical networks and random walks on undirected graphs. Choose two vertices a and b . For reference purposes let the voltage v_b equal zero. Attach a current source between a and b so that the voltage v_a equals one. Fixing the voltages at v_a and v_b induces voltages at all other vertices along with a current flow through the edges of the network. The analogy between electrical networks and random walks is the following. Having fixed the voltages at the vertices a and b , the voltage at an arbitrary vertex x equals the probability of a random walk starting at x reaching a before reaching b . If the voltage v_a is adjusted so that the current flowing into vertex a corresponds to one walk, then the current flowing through an edge is the net frequency with which a random walk from a to b traverses the edge.

Probabilistic interpretation of voltages

Before showing that the voltage at an arbitrary vertex x equals the probability of a random walk starting at x reaching a before reaching b , we first show that the voltages form a harmonic function. Let x and y be adjacent vertices and let i_{xy} be the current flowing through the edge from x to y . By Ohm's law,

$$i_{xy} = \frac{v_x - v_y}{r_{xy}} = (v_x - v_y)c_{xy}.$$

By Kirchoff's Law the currents flowing out of each vertex sum to zero.

$$\sum_y i_{xy} = 0$$

Replacing currents in the above sum by the voltage difference times the conductance yields

$$\sum_y (v_x - v_y)c_{xy} = 0$$

or

$$v_x \sum_y c_{xy} = \sum_y v_y c_{xy}.$$

Observing that $\sum_y c_{xy} = c_x$ and that $p_{xy} = \frac{c_{xy}}{c_x}$, yields $v_x c_x = \sum_y v_y p_{xy} c_x$. Hence, $v_x = \sum_y v_y p_{xy}$. Thus, the voltage at each vertex x is a weighted average of the voltages at the adjacent vertices. Hence the voltages form a harmonic function with $\{a, b\}$ as the boundary.

Let p_x be the probability that a random walk starting at vertex x reaches a before b . Clearly $p_a = 1$ and $p_b = 0$. Since $v_a = 1$ and $v_b = 0$, it follows that $p_a = v_a$ and $p_b = v_b$.

Furthermore, the probability of the walk reaching a from x before reaching b is the sum over all y adjacent to x of the probability of the walk going from x to y in the first step and then reaching a from y before reaching b . That is

$$p_x = \sum_y p_{xy} p_y.$$

Hence, p_x is the same harmonic function as the voltage function v_x and $\mathbf{v}$ and $\mathbf{p}$ satisfy the same boundary conditions at a and b . Thus, they are identical functions. The probability of a walk starting at x reaching a before reaching b is the voltage v_x .

Probabilistic interpretation of current

In a moment, we will set the current into the network at a to have a value which we will equate with one random walk. We will then show that the current i_{xy} is the net frequency with which a random walk from a to b goes through the edge xy before reaching b . Let u_x be the expected number of visits to vertex x on a walk from a to b before reaching b . Clearly $u_b = 0$. Every time the walk visits x , x not equal to a , it must come to x from some vertex y . Thus, the number of visits to x before reaching b is the sum over all y of the number of visits u_y to y before reaching b times the probability p_{yx} of going from y to x . For x not equal to b or a

$$u_x = \sum_{y \neq b} u_y p_{yx}.$$

Since $u_b = 0$ and $c_x p_{xy} = c_y p_{yx}$

$$u_x = \sum_{\text{all } y} u_y \frac{c_x p_{xy}}{c_y}$$

and hence $\frac{u_x}{c_x} = \sum_y \frac{u_y}{c_y} p_{xy}$. It follows that $\frac{u_x}{c_x}$ is harmonic with a and b as the boundary where the boundary conditions are $u_b = 0$ and u_a equals some fixed value. Now, $\frac{u_b}{c_b} = 0$. Setting the current into a to one, fixed the value of v_a . Adjust the current into a so that v_a equals $\frac{u_a}{c_a}$. Now $\frac{u_x}{c_x}$ and v_x satisfy the same harmonic conditions and thus are the same harmonic function. Let the current into a correspond to one walk. Note that if the walk starts at a and ends at b , the expected value of the difference between the number of times the walk leaves a and enters a must be one. This implies that the amount of current into a corresponds to one walk.

Next we need to show that the current i_{xy} is the net frequency with which a random walk traverses edge xy .

$$i_{xy} = (v_x - v_y) c_{xy} = \left(\frac{u_x}{c_x} - \frac{u_y}{c_y} \right) c_{xy} = u_x \frac{c_{xy}}{c_x} - u_y \frac{c_{xy}}{c_y} = u_x p_{xy} - u_y p_{yx}$$

The quantity $u_x p_{xy}$ is the expected number of times the edge xy is traversed from x to y and the quantity $u_y p_{yx}$ is the expected number of times the edge xy is traversed from y to

x . Thus, the current i_{xy} is the expected net number of traversals of the edge xy from x to y .

Effective resistance and escape probability

Set $v_a = 1$ and $v_b = 0$. Let i_a be the current flowing into the network at vertex a and out at vertex b . Define the *effective resistance* r_{eff} between a and b to be $r_{eff} = \frac{v_a}{i_a}$ and the *effective conductance* c_{eff} to be $c_{eff} = \frac{1}{r_{eff}}$. Define the *escape probability*, p_{escape} , to be the probability that a random walk starting at a reaches b before returning to a . We now show that the escape probability is $\frac{c_{eff}}{c_a}$. For convenience, assume that a and b are not adjacent. A slight modification of our argument suffices for the case when a and b are adjacent.

$$i_a = \sum_y (v_a - v_y) c_{ay}$$

Since $v_a = 1$,

$$\begin{aligned} i_a &= \sum_y c_{ay} - c_a \sum_y v_y \frac{c_{ay}}{c_a} \\ &= c_a \left[1 - \sum_y p_{ay} v_y \right]. \end{aligned}$$

For each y adjacent to the vertex a , p_{ay} is the probability of the walk going from vertex a to vertex y . Earlier we showed that v_y is the probability of a walk starting at y going to a before reaching b . Thus, $\sum_y p_{ay} v_y$ is the probability of a walk starting at a returning to a before reaching b and $1 - \sum_y p_{ay} v_y$ is the probability of a walk starting at a reaching b before returning to a . Thus, $i_a = c_a p_{escape}$. Since $v_a = 1$ and $c_{eff} = \frac{i_a}{v_a}$, it follows that $c_{eff} = i_a$. Thus, $c_{eff} = c_a p_{escape}$ and hence $p_{escape} = \frac{c_{eff}}{c_a}$.

For a finite connected graph the escape probability will always be nonzero. Now consider an infinite graph such as a lattice and a random walk starting at some vertex a . Form a series of finite graphs by merging all vertices at distance d or greater from a into a single vertex b for larger and larger values of d . The limit of p_{escape} as d goes to infinity is the probability that the random walk will never return to a . If $p_{escape} \rightarrow 0$, then eventually any random walk will return to a . If $p_{escape} \rightarrow q$ where $q > 0$, then a fraction of the walks never return. Thus, the escape probability terminology.

5.3 Random Walks on Undirected Graphs with Unit Edge Weights

We now focus our discussion on random walks on undirected graphs with uniform edge weights. At each vertex, the random walk is equally likely to take any edge. This corresponds to an electrical network in which all edge resistances are one. Assume the graph is connected. We consider questions such as what is the expected time for a random

walk starting at a vertex x to reach a target vertex y , what is the expected time until the random walk returns to the vertex it started at, and what is the expected time to reach every vertex?

Hitting time

The *hitting time* h_{xy} , sometimes called *discovery time*, is the expected time of a random walk starting at vertex x to reach vertex y . Sometimes a more general definition is given where the hitting time is the expected time to reach a vertex y from a given starting probability distribution.

One interesting fact is that adding edges to a graph may either increase or decrease h_{xy} depending on the particular situation. Adding an edge can shorten the distance from x to y thereby decreasing h_{xy} or the edge could increase the probability of a random walk going to some far off portion of the graph thereby increasing h_{xy} . Another interesting fact is that hitting time is not symmetric. The expected time to reach a vertex y from a vertex x in an undirected graph may be radically different from the time to reach x from y .

We start with two technical lemmas. The first lemma states that the expected time to traverse a path of n vertices is $\Theta(n^2)$.

Lemma 5.4 *The expected time for a random walk starting at one end of a path of n vertices to reach the other end is $\Theta(n^2)$.*

Proof: Consider walking from vertex 1 to vertex n in a graph consisting of a single path of n vertices. Let h_{ij} , $i < j$, be the hitting time of reaching j starting from i . Now $h_{12} = 1$ and

$$h_{i,i+1} = \frac{1}{2} + \frac{1}{2}(1 + h_{i-1,i+1}) = 1 + \frac{1}{2}(h_{i-1,i} + h_{i,i+1}) \quad 2 \leq i \leq n-1.$$

Solving for $h_{i,i+1}$ yields the recurrence

$$h_{i,i+1} = 2 + h_{i-1,i}.$$

Solving the recurrence yields

$$h_{i,i+1} = 2i - 1.$$

To get from 1 to n , go from 1 to 2, 2 to 3, etc. Thus

$$\begin{aligned} h_{1,n} &= \sum_{i=1}^{n-1} h_{i,i+1} = \sum_{i=1}^{n-1} (2i - 1) \\ &= 2 \sum_{i=1}^{n-1} i - \sum_{i=1}^{n-1} 1 \\ &= 2 \frac{n(n-1)}{2} - (n-1) \\ &= (n-1)^2. \end{aligned}$$

■

The lemma says that in a random walk on a line where we are equally likely to take one step to the right or left each time, the farthest we will go away from the start in n steps is $\Theta(\sqrt{n})$.

The next lemma shows that the expected time spent at vertex i by a random walk from vertex 1 to vertex n in a chain of n vertices is $2(i-1)$ for $2 \leq i \leq n-1$.

Lemma 5.5 *Consider a random walk from vertex 1 to vertex n in a chain of n vertices. Let $t(i)$ be the expected time spent at vertex i . Then*

$$t(i) = \begin{cases} n-1 & i=1 \\ 2(n-i) & 2 \leq i \leq n-1 \\ 1 & i=n. \end{cases}$$

Proof: Now $t(n) = 1$ since the walk stops when it reaches vertex n . Half of the time when the walk is at vertex $n-1$ it goes to vertex n . Thus $t(n-1) = 2$. For $3 \leq i < n-1$, $t(i) = \frac{1}{2}[t(i-1) + t(i+1)]$ and $t(1)$ and $t(2)$ satisfy $t(1) = \frac{1}{2}t(2) + 1$ and $t(2) = t(1) + \frac{1}{2}t(3)$. Solving for $t(i+1)$ for $3 \leq i < n-1$ yields

$$t(i+1) = 2t(i) - t(i-1)$$

which has solution $t(i) = 2(n-i)$ for $3 \leq i < n-1$. Then solving for $t(2)$ and $t(1)$ yields $t(2) = 2(n-2)$ and $t(1) = n-1$. Thus, the total time spent at vertices is

$$n-1 + 2(1+2+\cdots+n-2) + 1 = (n-1) + 2\frac{(n-1)(n-2)}{2} + 1 = (n-1)^2 + 1$$

which is one more than h_{1n} and thus is correct. ■

Adding edges to a graph might either increase or decrease the hitting time h_{xy} . Consider the graph consisting of a single path of n vertices. Add edges to this graph to get the graph in Figure 5.3 consisting of a clique of size $n/2$ connected to a path of $n/2$ vertices. Then add still more edges to get a clique of size n . Let x be the vertex at the midpoint of the original path and let y be the other endpoint of the path consisting of $n/2$ vertices as shown in the figure. In the first graph consisting of a single path of length n , $h_{xy} = \Theta(n^2)$. In the second graph consisting of a clique of size $n/2$ along with a path of length $n/2$, $h_{xy} = \Theta(n^3)$. To see this latter statement, note that starting at x , the walk will go down the path towards y and return to x $n/2$ times on average before reaching y for the first time. Each time the walk in the path returns to x , with probability $(n/2-1)/(n/2)$ it enters the clique and thus on average enters the clique $\Theta(n)$ times before starting down the path again. Each time it enters the clique, it spends $\Theta(n)$ time in the clique before returning to x . Thus, each time the walk returns to x from the path it spends $\Theta(n^2)$ time in the clique before starting down the path towards y for a total expected time that is

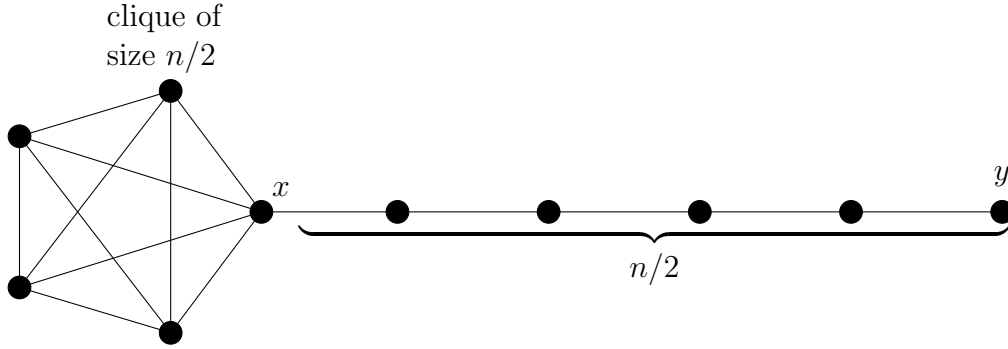


Figure 5.3: Illustration that adding edges to a graph can either increase or decrease hitting time.

$\Theta(n^3)$ before reaching y . In the third graph, which is the clique of size n , $h_{xy} = \Theta(n)$. Thus, adding edges first increased h_{xy} from n^2 to n^3 and then decreased it to n .

Hitting time is not symmetric even in the case of undirected graphs. In the graph of Figure 5.3, the expected time, h_{xy} , of a random walk from x to y , where x is the vertex of attachment and y is the other end vertex of the chain, is $\Theta(n^3)$. However, h_{yx} is $\Theta(n^2)$.

Commute time

The *commute time*, $\text{commute}(x, y)$, is the expected time of a random walk starting at x reaching y and then returning to x . So $\text{commute}(x, y) = h_{xy} + h_{yx}$. Think of going from home to office and returning home. We now relate the commute time to an electrical quantity, the effective resistance. The *effective resistance* between two vertices x and y in an electrical network is the voltage difference between x and y when one unit of current is inserted at vertex x and withdrawn from vertex y .

Theorem 5.6 *Given an undirected graph, consider the electrical network where each edge of the graph is replaced by a one ohm resistor. Given vertices x and y , the commute time, $\text{commute}(x, y)$, equals $2mr_{xy}$ where r_{xy} is the effective resistance from x to y and m is the number of edges in the graph.*

Proof: Insert at each vertex i a current equal to the degree d_i of vertex i . The total current inserted is $2m$ where m is the number of edges. Extract from a specific vertex j all of this $2m$ current. Let v_{ij} be the voltage difference from i to j . The current into i divides into the d_i resistors at vertex i . The current in each resistor is proportional to the voltage across it. Let k be a vertex adjacent to i . Then the current through the resistor between i and k is $v_{ij} - v_{kj}$, the voltage drop across the resistor. The sum of the currents out of i through the resistors must equal d_i , the current injected into i .

$$d_i = \sum_{\substack{k \text{ adj} \\ \text{to } i}} (v_{ij} - v_{kj}) = d_i v_{ij} - \sum_{\substack{k \text{ adj} \\ \text{to } i}} v_{kj}.$$

Solving for v_{ij}

$$v_{ij} = 1 + \sum_{\substack{k \text{ adj} \\ \text{to } i}} \frac{1}{d_i} v_{kj} = \sum_{\substack{k \text{ adj} \\ \text{to } i}} \frac{1}{d_i} (1 + v_{kj}). \quad (5.1)$$

Now the hitting time from i to j is the average time over all paths from i to k adjacent to i and then on from k to j . This is given by

$$h_{ij} = \sum_{\substack{k \text{ adj} \\ \text{to } i}} \frac{1}{d_i} (1 + h_{kj}). \quad (5.2)$$

Subtracting (5.2) from (5.1), gives $v_{ij} - h_{ij} = \sum_{\substack{k \text{ adj} \\ \text{to } i}} \frac{1}{d_i} (v_{kj} - h_{kj})$. Thus, the function $v_{ij} - h_{ij}$ is harmonic. Designate vertex j as the only boundary vertex. The value of $v_{ij} - h_{ij}$ at $i = j$, namely $v_{jj} - h_{jj}$, is zero, since both v_{jj} and h_{jj} are zero. So the function $v_{ij} - h_{ij}$ must be zero everywhere. Thus, the voltage v_{ij} equals the expected time h_{ij} from i to j .

To complete the proof, note that $h_{ij} = v_{ij}$ is the voltage from i to j when currents are inserted at all vertices in the graph and extracted at vertex j . If the current is extracted from i instead of j , then the voltages change and $v_{ji} = h_{ji}$ in the new setup. Finally, reverse all currents in this latter step. The voltages change again and for the new voltages $-v_{ji} = h_{ji}$. Since $-v_{ji} = v_{ij}$, we get $h_{ji} = v_{ij}$.

Thus, when a current is inserted at each vertex equal to the degree of the vertex and the current is extracted from j , the voltage v_{ij} in this set up equals h_{ij} . When we extract the current from i instead of j and then reverse all currents, the voltage v_{ij} in this new set up equals h_{ji} . Now, superpose both situations, i.e., add all the currents and voltages. By linearity, for the resulting v_{ij} , which is the sum of the other two v_{ij} 's, is $v_{ij} = h_{ij} + h_{ji}$. All currents cancel except the $2m$ amps injected at i and withdrawn at j . Thus, $2mr_{ij} = v_{ij} = h_{ij} + h_{ji} = \text{commute}(i, j)$ or $\text{commute}(i, j) = 2mr_{ij}$ where r_{ij} is the effective resistance from i to j . ■

The following corollary follows from Theorem 5.6 since the effective resistance r_{uv} is less than or equal to one when u and v are connected by an edge.

Corollary 5.7 *If vertices x and y are connected by an edge, then $h_{xy} + h_{yx} \leq 2m$ where m is the number of edges in the graph.*

Proof: If x and y are connected by an edge, then the effective resistance r_{xy} is less than or equal to one. ■

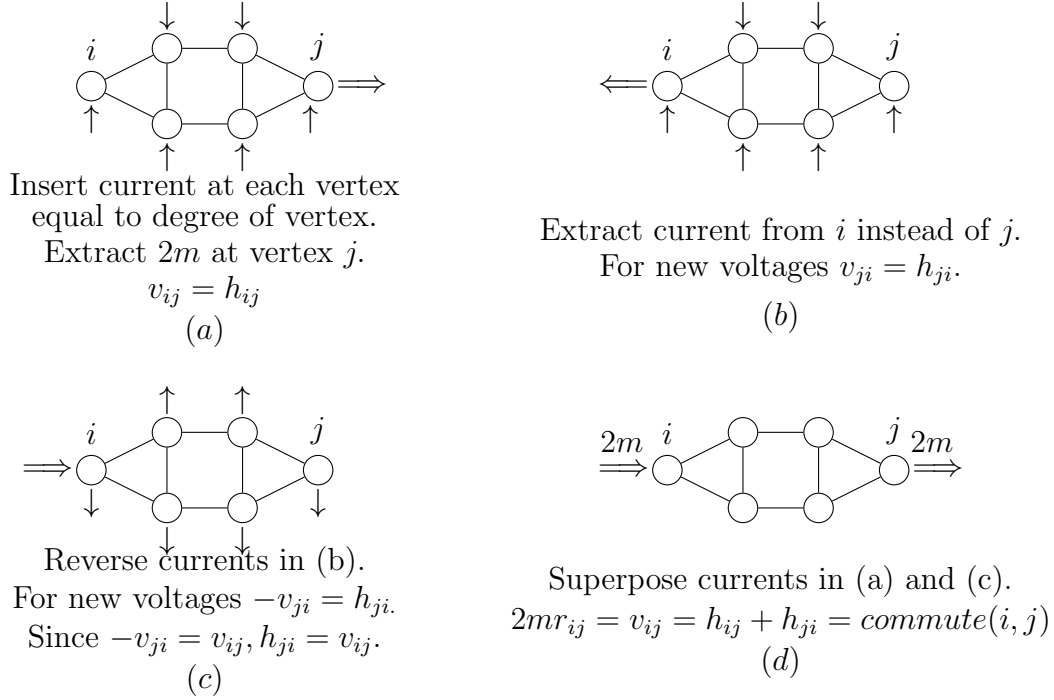


Figure 5.4: Illustration of proof that $\text{commute}(x, y) = 2mr_{xy}$ where m is the number of edges in the undirected graph and r_{xy} is the effective resistance between x and y .

Corollary 5.8 For vertices x and y in an n vertex graph, the commute time, $\text{commute}(x, y)$, is less than or equal to n^3 .

Proof: By Theorem 5.6 the commute time is given by the formula $\text{commute}(x, y) = 2mr_{xy}$ where m is the number of edges. In an n vertex graph there exists a path from x to y of length at most n . This implies $r_{xy} \leq n$ since the resistance can not be greater than that of any path from x to y . Since the number of edges is at most $\binom{n}{2}$

$$\text{commute}(x, y) = 2mr_{xy} \leq 2\binom{n}{2}n \cong n^3.$$

■

Again adding edges to a graph may increase or decrease the commute time. To see this, consider the graph consisting of a chain of n vertices, the graph of Figure 5.3, and the clique on n vertices.

Cover time

The *cover time*, $\text{cover}(x, G)$, is the expected time of a random walk starting at vertex x in the graph G to reach each vertex at least once. We write $\text{cover}(x)$ when G is understood.

The cover time of an undirected graph G , denoted $\text{cover}(G)$, is

$$\text{cover}(G) = \max_x \text{cover}(x, G).$$

For cover time of an undirected graph, increasing the number of edges in the graph may increase or decrease the cover time depending on the situation. Again consider three graphs, a chain of length n which has cover time $\Theta(n^2)$, the graph in Figure 5.3 which has cover time $\Theta(n^3)$, and the complete graph on n vertices which has cover time $\Theta(n \log n)$. Adding edges to the chain of length n to create the graph in Figure 5.3 increases the cover time from n^2 to n^3 and then adding even more edges to obtain the complete graph reduces the cover time to $n \log n$.

Note: The cover time of a clique is $\theta(n \log n)$ since this is the time to select every integer out of n integers with high probability, drawing integers at random. This is called the *coupon collector problem*. The cover time for a straight line is $\Theta(n^2)$ since it is the same as the hitting time. For the graph in Figure 5.3, the cover time is $\Theta(n^3)$ since one takes the maximum over all start states and $\text{cover}(x, G) = \Theta(n^3)$ where x is the vertex of attachment.

Theorem 5.9 *Let G be a connected graph with n vertices and m edges. The time for a random walk to cover all vertices of the graph G is bounded above by $4m(n - 1)$.*

Proof: Consider a depth first search of the graph G starting from some vertex z and let T be the resulting depth first search spanning tree of G . The depth first search covers every vertex. Consider the expected time to cover every vertex in the order visited by the depth first search. Clearly this bounds the cover time of G starting from vertex z . Note that each edge in T is traversed twice, once in each direction.

$$\text{cover}(z, G) \leq \sum_{\substack{(x,y) \in T \\ (y,x) \in T}} h_{xy}.$$

If (x, y) is an edge in T , then x and y are adjacent and thus Corollary 5.7 implies $h_{xy} \leq 2m$. Since there are $n - 1$ edges in the dfs tree and each edge is traversed twice, once in each direction, $\text{cover}(z) \leq 4m(n - 1)$. This holds for all starting vertices z . Thus, $\text{cover}(G) \leq 4m(n - 1)$ ■

The theorem gives the correct answer of n^3 for the $n/2$ clique with the $n/2$ tail. It gives an upper bound of n^3 for the n -clique where the actual cover time is $n \log n$.

Let r_{xy} be the effective resistance from x to y . Define the resistance $r_{eff}(G)$ of a graph G by $r_{eff}(G) = \max_{x,y} (r_{xy})$.

Theorem 5.10 *Let G be an undirected graph with m edges. Then the cover time for G is bounded by the following inequality*

$$mr_{eff}(G) \leq \text{cover}(G) \leq 2e^3 mr_{eff}(G) \ln n + n$$

where $e=2.71$ is Euler's constant and $r_{eff}(G)$ is the resistance of G .

Proof: By definition $r_{eff}(G) = \max_{x,y} (r_{xy})$. Let u and v be the vertices of G for which r_{xy} is maximum. Then $r_{eff}(G) = r_{uv}$. By Theorem 5.6, $\text{commute}(u, v) = 2mr_{uv}$. Hence $mr_{uv} = \frac{1}{2}\text{commute}(u, v)$. Clearly the commute time from u to v and back to u is less than twice the $\max(h_{uv}, h_{vu})$ and $\max(h_{uv}, h_{vu})$ is clearly less than the cover time of G . Putting these facts together gives the first inequality in the theorem.

$$mr_{eff}(G) = mr_{uv} = \frac{1}{2}\text{commute}(u, v) \leq \max(h_{uv}, h_{vu}) \leq \text{cover}(G)$$

For the second inequality in the theorem, by Theorem 5.6, for any x and y , $\text{commute}(x, y)$ equals $2mr_{xy}$ which is less than or equal to $2mr_{eff}(G)$, implying $h_{xy} \leq 2mr_{eff}(G)$. By the Markov inequality, since the expected time to reach y starting at any x is less than $2mr_{eff}(G)$, the probability that y is not reached from x in $2mr_{eff}(G)e^3$ steps is at most $\frac{1}{e^3}$. Thus, the probability that a vertex y has not been reached in $2e^3 mr_{eff}(G) \log n$ steps is at most $\frac{1}{e^3} \ln n = \frac{1}{n^3}$ because a random walk of length $2e^3 mr(G) \log n$ is a sequence of $\log n$ independent random walks, each of length $2e^3 mr(G) r_{eff}(G)$. Suppose after a walk of $2e^3 mr_{eff}(G) \log n$ steps, vertices $v_1, v_2, \dots, v_l$ had not been reached. Walk until v_1 is reached, then v_2 , etc. By Corollary 5.8 the expected time for each of these is n^3 , but since each happens only with probability $1/n^3$, we effectively take $O(1)$ time per v_i , for a total time at most n . More precisely,

$$\begin{aligned} \text{cover}(G) &\leq 2e^3 mr_{eff}(G) + \sum_v \text{Prob}(v \text{ was not visited in the first } 2e^3 mr_{eff}(G) \text{ steps}) n^3 \\ &\leq 2e^3 mr_{eff}(G) + \sum_v \frac{1}{n^3} n^3 \leq 2e^3 mr_{eff}(G) + n. \end{aligned}$$

■

5.4 Random Walks in Euclidean Space

Many physical processes such as Brownian motion are modeled by random walks. Random walks in Euclidean d -space consisting of fixed length steps parallel to the coordinate axes are really random walks on a d -dimensional lattice and are a special case of random walks on graphs. In a random walk on a graph, at each time unit an edge from the current vertex is selected at random and the walk proceeds to the adjacent vertex. We begin by studying random walks on lattices.

Random walks on lattices

We now apply the analogy between random walks and current to lattices. Consider a random walk on a finite segment $-n, \dots, -1, 0, 1, 2, \dots, n$ of a one dimensional lattice starting from the origin. Is the walk certain to return to the origin or is there some probability that it will escape, i.e., reach the boundary before returning? The probability of reaching the boundary before returning to the origin is called the escape probability. We shall be interested in this quantity as n goes to infinity.

Convert the lattice to an electrical network by replacing each edge with a one ohm resistor. Then the probability of a walk starting at the origin reaching n or $-n$ before returning to the origin is the escape probability given by

$$p_{\text{escape}} = \frac{c_{\text{eff}}}{c_a}$$

where c_{eff} is the effective conductance between the origin and the boundary points and c_a is the sum of the conductance's at the origin. In a d -dimensional lattice, $c_a = 2d$ assuming that the resistors have value one. For the d -dimensional lattice

$$p_{\text{escape}} = \frac{1}{2d r_{\text{eff}}}$$

In one dimension, the electrical network is just two series connections of n one ohm resistors connected in parallel. So as n goes to infinity, r_{eff} goes to infinity and the escape probability goes to zero as n goes to infinity. Thus, the walk in the unbounded one dimensional lattice will return to the origin with probability one. This is equivalent to flipping a balanced coin and keeping track of the number of heads minus the number of tails. The count will return to zero infinitely often. By the law of large numbers in n steps with high probability the walk will be within $\sqrt{n}$ distance of the origin.

Two dimensions

For the 2-dimensional lattice, consider a larger and larger square about the origin for the boundary as shown in Figure 5.5a and consider the limit of r_{eff} as the squares get larger. Shorting the resistors on each square can only reduce r_{eff} . Shorting the resistors results in the linear network shown in Figure 5.5b. As the paths get longer, the number of resistors in parallel also increases. The resistor between vertex i and $i + 1$ is really $4(2i + 1)$ unit resistors in parallel. The effective resistance of $4(2i + 1)$ resistors in parallel is $1/4(2i + 1)$. Thus,

$$r_{\text{eff}} \geq \frac{1}{4} + \frac{1}{12} + \frac{1}{20} + \dots = \frac{1}{4}(1 + \frac{1}{3} + \frac{1}{5} + \dots) = \Theta(\ln n).$$

Since the lower bound on the effective resistance and hence the effective resistance goes to infinity, the escape probability goes to zero for the 2-dimensional lattice.

Three dimensions

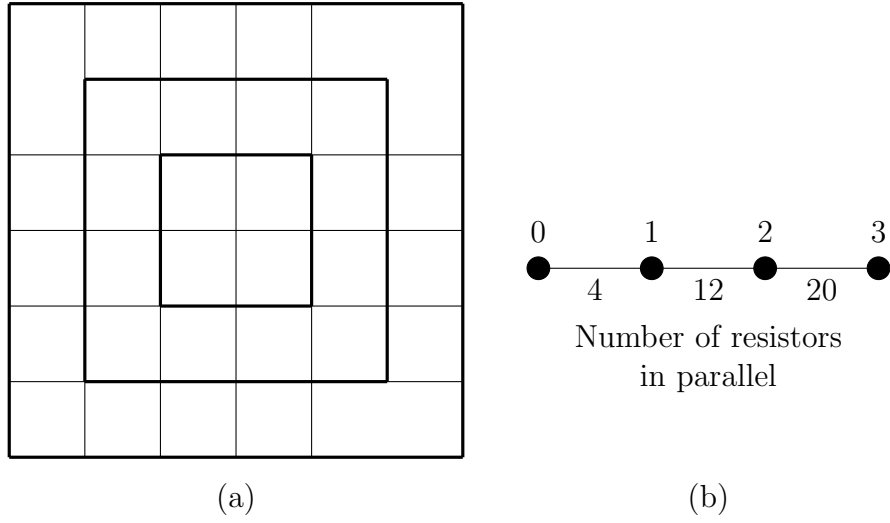


Figure 5.5: 2-dimensional lattice along with the linear network resulting from shorting resistors on the concentric squares about the origin.

In three dimensions, the resistance along any path to infinity grows to infinity but the number of paths in parallel also grows to infinity. It turns out there are a sufficient number of paths that r_{eff} remains finite and thus there is a nonzero escape probability. We will prove this now. First note that shorting any edge decreases the resistance, so we do not use shorting in this proof, since we seek to prove an upper bound on the resistance. Instead we remove some edges, which increases their resistance to infinity and hence increases the effective resistance, giving an upper bound. To simplify things we consider walks on one quadrant rather than the full grid. The resistance to infinity derived from only the quadrant is an upper bound on the resistance of the full grid.

The construction used in three dimensions is easier to explain first in two dimensions. Draw dotted diagonal lines at $x + y = 2^n - 1$. Consider two paths that start at the origin. One goes up and the other goes to the right. Each time a path encounters a dotted diagonal line, split the path into two, one which goes right and the other up. Where two paths cross, split the vertex into two, keeping the paths separate. By a symmetry argument, splitting the vertex does not change the resistance of the network. Remove all resistors except those on these paths. The resistance of the original network is less than that of the tree produced by this process since removing a resistor is equivalent to increasing its resistance to infinity.

The distances between splits increase and are 1, 2, 4, etc. At each split the number of paths in parallel doubles. See Figure 5.7. Thus, the resistance to infinity in this two dimensional example is

$$\frac{1}{2} + \frac{1}{4}2 + \frac{1}{8}4 + \cdots = \frac{1}{2} + \frac{1}{2} + \frac{1}{2} + \cdots = \infty.$$

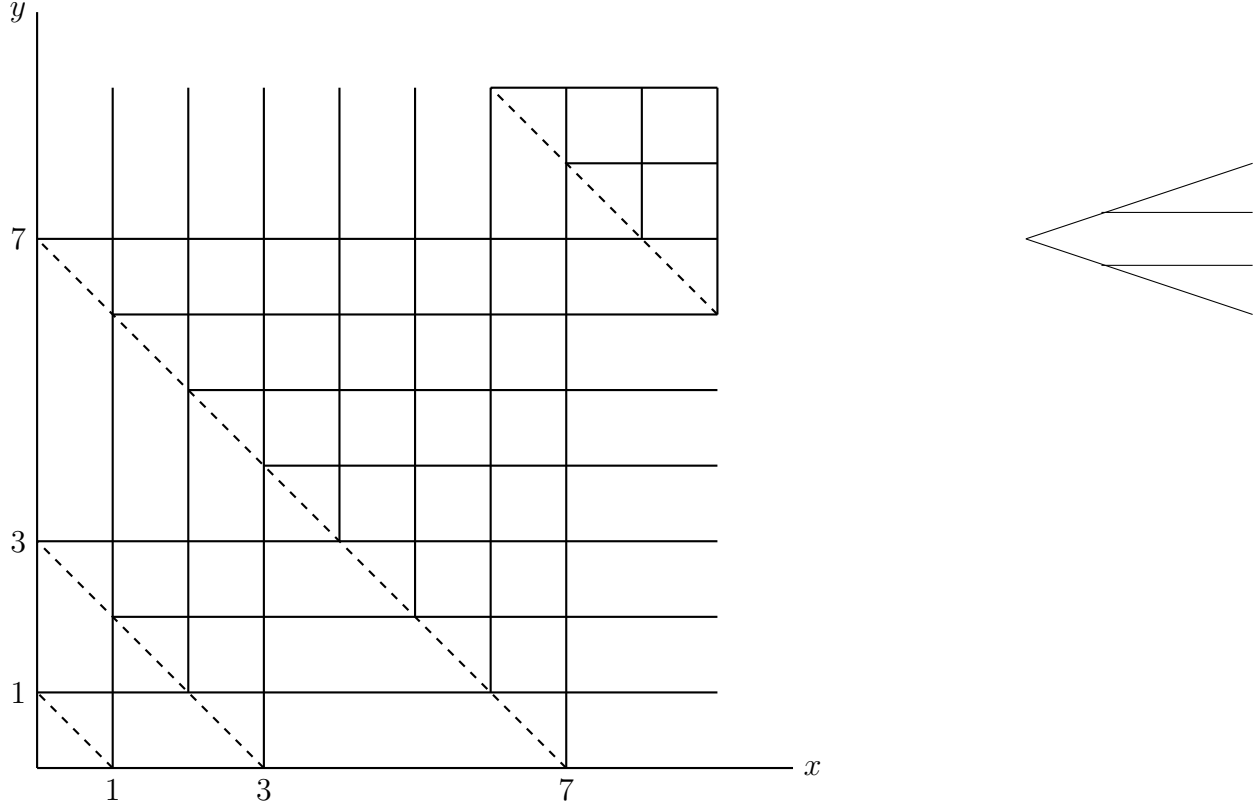


Figure 5.6: Paths in a 2-dimensional lattice obtained from the 3-dimensional construction applied in 2-dimensions.

In the analogous three dimensional construction, paths go up, to the right, and out of the plane of the paper. The paths split three ways at planes given by $x + y + z = 2^n - 1$. Each time the paths split the number of parallel segments triple. Segments of the paths between splits are of length 1, 2, 4, etc. and the resistance of the segments are equal to the lengths. The resistance out to infinity for the tree is

$$\frac{1}{3} + \frac{1}{9}2 + \frac{1}{27}4 + \cdots = \frac{1}{3} \left(1 + \frac{2}{3} + \frac{4}{9} + \cdots \right) = \frac{1}{3} \frac{1}{1 - \frac{2}{3}} = 1$$

The resistance of the three dimensional lattice is less. It is important to check that the paths are edge-disjoint and so the tree is a subgraph of the lattice. Going to a subgraph is equivalent to deleting edges which only increases the resistance. That is why the resistance of the lattice is less than that of the tree. Thus, in three dimensions the escape probability is nonzero. The upper bound on r_{eff} gives the lower bound

$$p_{escape} = \frac{1}{2d} \frac{1}{r_{eff}} \geq \frac{1}{6}.$$

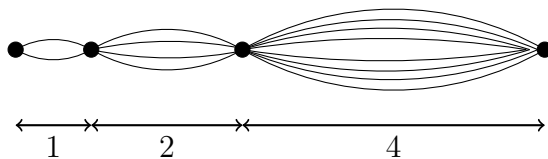


Figure 5.7: Paths obtained from 2-dimensional lattice. Distances between splits double as do the number of parallel paths.

A lower bound on r_{eff} gives an upper bound on p_{escape} . To get the upper bound on p_{escape} , short all resistors on surfaces of boxes at distances 1, 2, 3, etc. Then

$$r_{eff} \geq \frac{1}{6} \left[1 + \frac{1}{9} + \frac{1}{25} + \cdots \right] \geq \frac{1.23}{6} \geq 0.2$$

This gives

$$p_{escape} = \frac{1}{2d} \frac{1}{r_{eff}} \leq \frac{5}{6}.$$

5.5 The Web as a Markov Chain

A modern application of random walks on directed graphs comes from trying to establish the importance of pages on the World Wide Web. One way to do this would be to take a random walk on the web viewed as a directed graph with an edge corresponding to each hypertext link and rank pages according to their stationary probability. A connected, undirected graph is strongly connected in that one can get from any vertex to any other vertex and back again. Often the directed case is not strongly connected. One difficulty occurs if there is a vertex with no out edges. When the walk encounters this vertex the walk disappears. Another difficulty is that a vertex or a strongly connected component with no in edges is never reached. One way to resolve these difficulties is to introduce a random restart condition. At each step, with some probability r , jump to a vertex selected uniformly at random and with probability $1 - r$ select an edge at random and follow it. If a vertex has no out edges, the value of r for that vertex is set to one. This has the effect of converting the graph to a strongly connected graph so that the stationary probabilities exist.

Page rank and hitting time

The page rank of a vertex in a directed graph is the stationary probability of the vertex, where we assume a positive restart probability of say $r = 0.15$. The restart ensures that the graph is strongly connected. The page rank of a page is the fractional frequency with which the page will be visited over a long period of time. If the page rank is p , then the expected time between visits or return time is $1/p$. Notice that one can increase the pagerank of a page by reducing the return time and this can be done by creating short cycles.

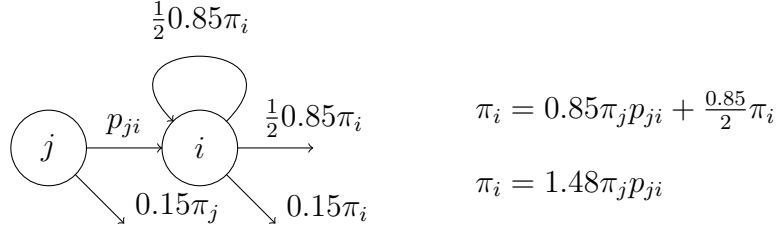


Figure 5.8: Impact on page rank of adding a self loop

Consider a vertex i with a single edge in from vertex j and a single edge out. The stationary probability π satisfies $\pi P = \pi$, and thus

$$\pi_i = \pi_j p_{ji}.$$

Adding a self-loop at i , results in a new equation

$$\pi_i = \pi_j p_{ji} + \frac{1}{2}\pi_i$$

or

$$\pi_i = 2 \pi_j p_{ji}.$$

Of course, π_j would have changed too, but ignoring this for now, pagerank is doubled by the addition of a self-loop. Adding k self loops, results in the equation

$$\pi_i = \pi_j p_{ji} + \frac{k}{k+1}\pi_i,$$

and again ignoring the change in π_j , we now have $\pi_i = (k+1)\pi_j p_{ji}$. What prevents one from increasing the page rank of a page arbitrarily? The answer is the restart. We neglected the 0.15 probability that is taken off for the random restart. With the restart taken into account, the equation for π_i when there is no self-loop is

$$\pi_i = 0.85\pi_j p_{ji}$$

whereas, with k self-loops, the equation is

$$\pi_i = 0.85\pi_j p_{ji} + 0.85\frac{k}{k+1}\pi_i.$$

Solving for π_i yields

$$\pi_i = \frac{0.85k + 0.85}{0.15k + 1}\pi_j p_{ji}$$

which for $k = 1$ is $\pi_i = 1.48\pi_j p_{ji}$ and in the limit as $k \rightarrow \infty$ is $\pi_i = 5.67\pi_j p_{ji}$. Adding a single loop only increases pagerank by a factor of 1.74 and adding k loops increases it by at most a factor of 6.67 for arbitrarily large k .

Hitting time

Related to page rank is a quantity called hitting time. Hitting time is closely related to return time and thus to the reciprocal of page rank. One way to return to a vertex v is by a path in the graph from v back to v . Another way is to start on a path that encounters a restart, followed by a path from the random restart vertex to v . The time to reach v after a restart is the hitting time. Thus, return time is clearly less than the expected time until a restart plus hitting time. The fastest one could return would be if there were only paths of length two since self loops are ignored in calculating page rank. If r is the restart value, then the loop would be traversed with at most probability $(1 - r)^2$. With probability $r + (1 - r)r = (2 - r)r$ one restarts and then hits v . Thus, the return time is at least $(1 - r)^2 + (2 - r)r \times (\text{hitting time})$. Combining these two bounds yields

$$(1 - r)^2 + (2 - r)rE(\text{hitting time}) \leq E(\text{return time}) \leq E(\text{hitting time}).$$

The relationship between return time and hitting time can be used to see if a vertex has unusually high probability of short loops. However, there is no efficient way to compute hitting time for all vertices as there is for return time. For a single vertex v , one can compute hitting time by removing the edges out of the vertex v for which one is computing hitting time and then run the page rank algorithm for the new graph. The hitting time for v is the reciprocal of the page rank in the graph with the edges out of v removed. Since computing hitting time for each vertex requires removal of a different set of edges, the algorithm only gives the hitting time for one vertex at a time. Since one is probably only interested in the hitting time of vertices with low hitting time, an alternative would be to use a random walk to estimate the hitting time of low hitting time vertices.

Spam

Suppose one has a web page and would like to increase its page rank by creating some other web pages with pointers to the original page. The abstract problem is the following. We are given a directed graph G and a vertex v whose page rank we want to increase. We may add new vertices to the graph and add edges from v or from the new vertices to any vertices we want. We cannot add edges out of other vertices. We can also delete edges from v .

The page rank of v is the stationary probability for vertex v with random restarts. If we delete all existing edges out of v , create a new vertex u and edges (v, u) and (u, v) , then the page rank will be increased since any time the random walk reaches v it will be captured in the loop $v \rightarrow u \rightarrow v$. A search engine can counter this strategy by more frequent random restarts.

A second method to increase page rank would be to create a star consisting of the vertex v at its center along with a large set of new vertices each with a directed edge to

v . These new vertices will sometimes be chosen as the target of the random restart and hence the vertices increase the probability of the random walk reaching v . This second method is countered by reducing the frequency of random restarts.

Notice that the first technique of capturing the random walk increases page rank but does not effect hitting time. One can negate the impact of someone capturing the random walk on page rank by increasing the frequency of random restarts. The second technique of creating a star increases page rank due to random restarts and decreases hitting time. One can check if the page rank is high and hitting time is low in which case the page rank is likely to have been artificially inflated by the page capturing the walk with short cycles.

Personalized page rank

In computing page rank, one uses a restart probability, typically 0.15, in which at each step, instead of taking a step in the graph, the walk goes to a vertex selected uniformly at random. In personalized page rank, instead of selecting a vertex uniformly at random, one selects a vertex according to a personalized probability distribution. Often the distribution has probability one for a single vertex and whenever the walk restarts it restarts at that vertex.

Algorithm for computing personalized page rank

First, consider the normal page rank. Let α be the restart probability with which the random walk jumps to an arbitrary vertex. With probability $1 - \alpha$ the random walk selects a vertex uniformly at random from the set of adjacent vertices. Let $\mathbf{p}$ be a row vector denoting the page rank and let G be the adjacency matrix with rows normalized to sum to one. Then

$$\mathbf{p} = \frac{\alpha}{n} (1, 1, \dots, 1) + (1 - \alpha) \mathbf{p} G$$

$$\mathbf{p}[I - (1 - \alpha)G] = \frac{\alpha}{n} (1, 1, \dots, 1)$$

or

$$\mathbf{p} = \frac{\alpha}{n} (1, 1, \dots, 1) [I - (1 - \alpha) G]^{-1}.$$

Thus, in principle, $\mathbf{p}$ can be found by computing the inverse of $[I - (1 - \alpha)G]^{-1}$. But this is far from practical since for the whole web one would be dealing with matrices with billions of rows and columns. A more practical procedure is to run the random walk and observe using the basics of the power method in Chapter 4 that the process converges to the solution $\mathbf{p}$.

For the personalized page rank, instead of restarting at an arbitrary vertex, the walk restarts at a designated vertex. More generally, it may restart in some specified neighborhood. Suppose the restart selects a vertex using the probability distribution s . Then, in

the above calculation replace the vector $\frac{1}{n}(1, 1, \dots, 1)$ by the vector $\mathbf{s}$. Again, the computation could be done by a random walk. But, we wish to do the random walk calculation for personalized pagerank quickly since it is to be performed repeatedly. With more care this can be done, though we do not describe it here.

5.6 Markov Chain Monte Carlo

The Markov Chain Monte Carlo (MCMC) method is a technique for sampling a multivariate probability distribution $p(\mathbf{x})$, where $\mathbf{x} = (x_1, x_2, \dots, x_d)$. The MCMC method is used to estimate the expected value of a function $f(\mathbf{x})$

$$E(f) = \sum_{\mathbf{x}} f(\mathbf{x})p(\mathbf{x}).$$

If each x_i can take on two or more values, then there are at least 2^d values for $\mathbf{x}$, so an explicit summation requires exponential time. Instead, one could draw a set of samples, each sample $\mathbf{x}$ with probability $p(\mathbf{x})$. Averaging f over these samples provides an estimate of the sum.

To sample according to $p(\mathbf{x})$, design a Markov Chain whose states correspond to the possible values of $\mathbf{x}$ and whose stationary probability distribution is $p(\mathbf{x})$. There are two general techniques to design such a Markov Chain: the Metropolis-Hastings algorithm and Gibbs sampling. The Fundamental Theorem of Markov Chains, Theorem 5.2, states that the average of f over states seen in a sufficiently long run is a good estimate of $E(f)$. The harder task is to show that the number of steps needed before the long-run average probabilities are close to the stationary distribution grows polynomially in d , though the total number of states may grow exponentially in d . This phenomenon known as *rapid mixing* happens for a number of interesting examples. Section 5.8 presents a crucial tool used to show rapid mixing.

We used $\mathbf{x} \in \mathbf{R}^d$ to emphasize that distributions are multi-variate. From a Markov chain perspective, each value $\mathbf{x}$ can take on is a state, i.e., a vertex of the graph on which the random walk takes place. Henceforth, we will use the subscripts $i, j, k, \dots$ to denote states and will use p_i instead of $p(x_1, x_2, \dots, x_d)$ to denote the probability of the state corresponding to a given set of values for the variables. Recall that in the Markov chain terminology, vertices of the graph are called states.

Recall the notation that $\mathbf{p}^{(t)}$ is the row vector of probabilities of the random walk being at each state (vertex of the graph) at time t . So, $\mathbf{p}^{(t)}$ has as many components as there are states and its i^{th} component, $p_i^{(t)}$, is the probability of being in state i at time t . Recall the long-term t -step average is

$$\mathbf{a}^{(t)} = \frac{1}{t} [\mathbf{p}^{(0)} + \mathbf{p}^{(1)} + \dots + \mathbf{p}^{(t-1)}]. \quad (5.3)$$

The expected value of the function f under the probability distribution $\mathbf{p}$ is $E(f) = \sum_i f_i p_i$ where f_i is the value of f at state i . Our estimate of this quantity will be the average value of f at the states seen in a t step walk. Call this estimate a . Clearly, the expected value of a is

$$E(a) = \sum_i f_i \left(\frac{1}{t} \sum_{j=1}^t \text{Prob}(\text{walk is in state } i \text{ at time } j) \right) = \sum_i f_i a_i^{(t)}.$$

The expectation here is with respect to the “coin tosses” of the algorithm, not with respect to the underlying distribution p . Let $f_{\max}$ denote the maximum absolute value of f . It is easy to see that

$$\left| \sum_i f_i p_i - E(a) \right| \leq f_{\max} \sum_i |p_i - a_i^{(t)}| = f_{\max} |\mathbf{p} - \mathbf{a}^{(t)}|_1 \quad (5.4)$$

where the quantity $|\mathbf{p} - \mathbf{a}^{(t)}|_1$ is the l_1 distance between the probability distributions $\mathbf{p}$ and $\mathbf{a}^{(t)}$ and is often called the “total variation distance” between the distributions. We will build tools to upper bound $|\mathbf{p} - \mathbf{a}^{(t)}|_1$. Since $\mathbf{p}$ is the stationary distribution, the t for which $|\mathbf{p} - \mathbf{a}^{(t)}|_1$ becomes small is determined by the rate of convergence of the Markov chain to its steady state.

The following proposition is often useful.

Proposition 5.11 *For two probability distributions $\mathbf{p}$ and $\mathbf{q}$,*

$$|\mathbf{p} - \mathbf{q}|_1 = 2 \sum_i (p_i - q_i)^+ = 2 \sum_i (q_i - p_i)^+$$

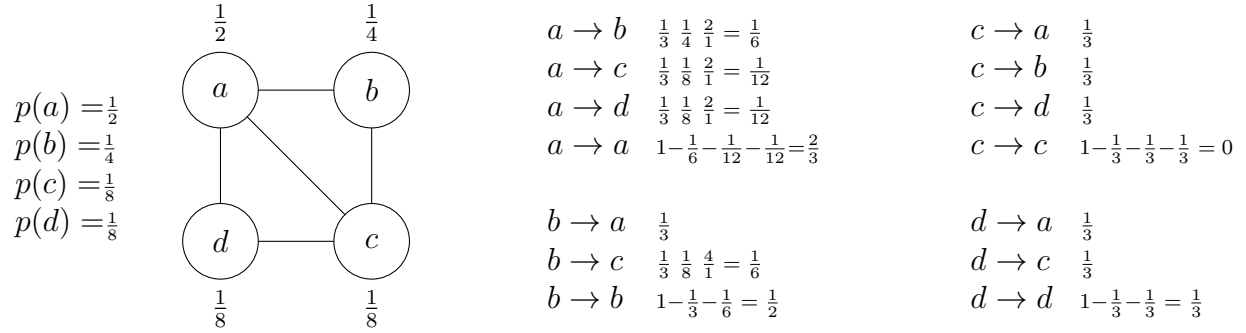
where $x^+ = x$ if $x \geq 0$ and $x^+ = 0$ if $x < 0$.

The proof is left as an exercise.

5.6.1 Metropolis-Hasting Algorithm

The Metropolis-Hasting algorithm is a general method to design a Markov chain whose stationary distribution is a given target distribution p . Start with a connected undirected graph G on the set of states. If the states are the lattice points $(x_1, x_2, \dots, x_d)$ in $\mathbf{R}^d$ with $x_i \in \{0, 1, 2, \dots, n\}$, then G is the lattice graph with $2d$ coordinate edges at each interior vertex. In general, let r be the maximum degree of any vertex of G . The transitions of the Markov chain are defined as follows. At state i select neighbor j with probability $\frac{1}{r}$. Since the degree of i may be less than r , with some probability no edge is selected and the walk remains at i . If a neighbor j is selected and $p_j \geq p_i$, go to j . If $p_j < p_i$, go to j with probability p_j/p_i and stay at i with probability $1 - \frac{p_j}{p_i}$. Intuitively, this favors “heavier” states with higher p values. So, for $i \neq j$, adjacent in G ,

$$p_{ij} = \frac{1}{r} \min \left(1, \frac{p_j}{p_i} \right)$$



$$p(a) = p(a)p(a \rightarrow a) + p(b)p(b \rightarrow a) + p(c)p(c \rightarrow a) + p(d)p(d \rightarrow a)$$

$$= \frac{1}{2} \cdot \frac{2}{3} + \frac{1}{4} \cdot \frac{1}{3} + \frac{1}{8} \cdot \frac{1}{3} + \frac{1}{8} \cdot \frac{1}{3} = \frac{1}{2}$$

$$p(b) = p(a)p(a \rightarrow b) + p(b)p(b \rightarrow b) + p(c)p(c \rightarrow b)$$

$$= \frac{1}{2} \cdot \frac{1}{6} + \frac{1}{4} \cdot \frac{1}{2} + \frac{1}{8} \cdot \frac{1}{3} = \frac{1}{4}$$

$$p(c) = p(a)p(a \rightarrow c) + p(b)p(b \rightarrow c) + p(c)p(c \rightarrow c) + p(d)p(d \rightarrow c)$$

$$= \frac{1}{2} \cdot \frac{1}{12} + \frac{1}{4} \cdot \frac{1}{6} + \frac{1}{8} \cdot 0 + \frac{1}{8} \cdot \frac{1}{3} = \frac{1}{8}$$

$$p(d) = p(a)p(a \rightarrow d) + p(c)p(c \rightarrow d) + p(d)p(d \rightarrow d)$$

$$= \frac{1}{2} \cdot \frac{1}{12} + \frac{1}{8} \cdot \frac{1}{3} + \frac{1}{8} \cdot \frac{1}{3} = \frac{1}{8}$$

Figure 5.9: Using the Metropolis-Hasting algorithm to set probabilities for a random walk so that the stationary probability will be the desired probability.

and

$$p_{ii} = 1 - \sum_{j \neq i} p_{ij}.$$

Thus,

$$p_i p_{ij} = \frac{p_i}{r} \min \left(1, \frac{p_j}{p_i} \right) = \frac{1}{r} \min(p_i, p_j) = \frac{p_j}{r} \min \left(1, \frac{p_i}{p_j} \right) = p_j p_{ji}.$$

By Lemma 5.3, the stationary probabilities are indeed $p(\mathbf{x})$ as desired.

Example: Consider the graph in Figure 5.9. Using the Metropolis-Hasting algorithm, assign transition probabilities so that the stationary probability of a random walk is $p(a) = \frac{1}{2}$, $p(b) = \frac{1}{4}$, $p(c) = \frac{1}{8}$, and $p(d) = \frac{1}{8}$. The maximum degree of any vertex is three, so at a , the probability of taking the edge (a, b) is $\frac{1}{3} \cdot \frac{1}{4} \cdot \frac{2}{1}$ or $\frac{1}{6}$. The probability of taking the edge (a, c) is $\frac{1}{3} \cdot \frac{1}{8} \cdot \frac{2}{1}$ or $\frac{1}{12}$ and of taking the edge (a, d) is $\frac{1}{3} \cdot \frac{1}{8} \cdot \frac{2}{1}$ or $\frac{1}{12}$. Thus, the probability of staying at a is $\frac{2}{3}$. The probability of taking the edge from b to a is $\frac{1}{3}$. The probability of taking the edge from c to a is $\frac{1}{3}$ and the probability of taking the edge from d to a is $\frac{1}{3}$. Thus, the stationary probability of a is $\frac{1}{4} \cdot \frac{1}{3} + \frac{1}{8} \cdot \frac{1}{3} + \frac{1}{8} \cdot \frac{1}{3} + \frac{1}{2} \cdot \frac{2}{3} = \frac{1}{2}$, which is the desired probability. ■

5.6.2 Gibbs Sampling

Gibbs sampling is another Markov Chain Monte Carlo method to sample from a multivariate probability distribution. Let $p(\mathbf{x})$ be the target distribution where $\mathbf{x} = (x_1, \dots, x_d)$. Gibbs sampling consists of a random walk on an undirected graph whose vertices correspond to the values of $\mathbf{x} = (x_1, \dots, x_d)$ and in which there is an edge from $\mathbf{x}$ to $\mathbf{y}$ if $\mathbf{x}$ and $\mathbf{y}$ differ in only one coordinate. Thus, the underlying graph is a d -dimensional lattice.

To generate samples of $\mathbf{x} = (x_1, \dots, x_d)$ with a target distribution $p(\mathbf{x})$, the Gibbs sampling algorithm repeats the following steps. One of the variables x_i is chosen to be updated. Its new value is chosen based on the marginal probability of x_i with the other variables fixed. There are two commonly used schemes to determine which x_i to update. One scheme is to choose x_i randomly, the other is to choose x_i by sequentially scanning from x_1 to x_d .

Suppose that $\mathbf{x}$ and $\mathbf{y}$ are two states that differ in only one coordinate. Without loss of generality let that coordinate be the first. Then, in the scheme where a coordinate is randomly chosen to modify, the probability $p_{\mathbf{xy}}$ of going from $\mathbf{x}$ to $\mathbf{y}$ is

$$p_{\mathbf{xy}} = \frac{1}{d} p(y_1 | x_2, x_3, \dots, x_d).$$

The normalizing constant is $1/d$ since for a given value i the probability distribution of $p(y_i | x_1, x_2, \dots, x_{i-1}, x_{i+1}, \dots, x_d)$ sums to one, and thus summing i over the d -dimensions results in a value of d . Similarly,

$$\begin{aligned} p_{\mathbf{yx}} &= \frac{1}{d} p(x_1 | y_2, y_3, \dots, y_d) \\ &= \frac{1}{d} p(x_1 | x_2, x_3, \dots, x_d). \end{aligned}$$

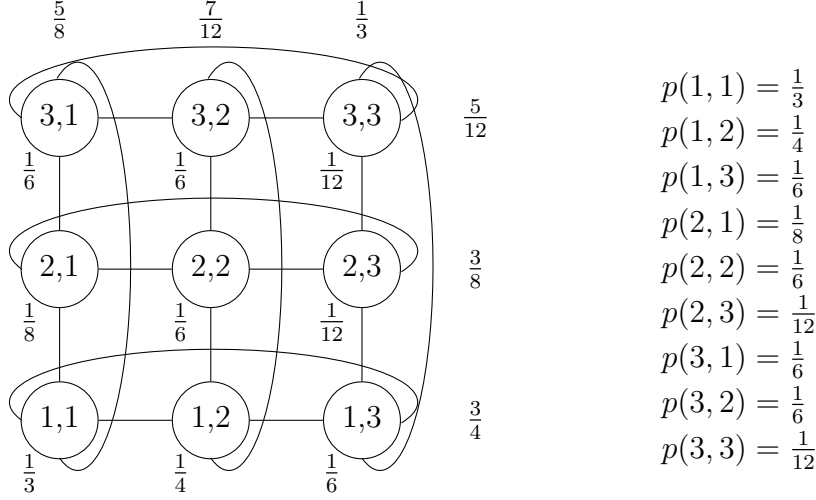
Here use was made of the fact that for $j \neq i$, $x_j = y_j$.

It is simple to see that this chain has stationary probability proportional to $p(\mathbf{x})$. Rewrite $p_{\mathbf{xy}}$ as

$$\begin{aligned} p_{\mathbf{xy}} &= \frac{1}{d} \frac{p(y_1 | x_2, x_3, \dots, x_d) p(x_2, x_3, \dots, x_d)}{p(x_2, x_3, \dots, x_d)} \\ &= \frac{1}{d} \frac{p(y_1, x_2, x_3, \dots, x_d)}{p(x_2, x_3, \dots, x_d)} \\ &= \frac{1}{d} \frac{p(\mathbf{y})}{p(x_2, x_3, \dots, x_d)} \end{aligned}$$

again using $x_j = y_j$ for $j \neq i$. Similarly write

$$p_{\mathbf{yx}} = \frac{1}{d} \frac{p(\mathbf{x})}{p(x_2, x_3, \dots, x_d)}$$



$$\begin{aligned}
p(1, 1) &= \frac{1}{3} \\
p(1, 2) &= \frac{1}{4} \\
p(1, 3) &= \frac{1}{6} \\
p(2, 1) &= \frac{1}{8} \\
p(2, 2) &= \frac{1}{6} \\
p(2, 3) &= \frac{1}{12} \\
p(3, 1) &= \frac{1}{6} \\
p(3, 2) &= \frac{1}{6} \\
p(3, 3) &= \frac{1}{12}
\end{aligned}$$

$$p_{(11)(12)} = \frac{1}{d} p_{12} / (p_{11} + p_{12} + p_{13}) = \frac{1}{2} \frac{1}{4} / (\frac{1}{3} \frac{1}{4} \frac{1}{6}) = \frac{1}{2} \frac{1}{4} / \frac{9}{12} = \frac{1}{2} \frac{1}{4} \frac{4}{3} = \frac{1}{6}$$

Calculation of edge probability $p_{(11)(12)}$

$$\begin{aligned}
p_{(11)(12)} &= \frac{1}{2} \frac{1}{4} \frac{4}{3} = \frac{1}{6} & p_{(12)(11)} &= \frac{1}{2} \frac{1}{3} \frac{4}{3} = \frac{2}{9} & p_{(13)(11)} &= \frac{1}{2} \frac{1}{3} \frac{4}{3} = \frac{2}{9} & p_{(21)(22)} &= \frac{1}{2} \frac{1}{6} \frac{8}{3} = \frac{2}{9} \\
p_{(11)(13)} &= \frac{1}{2} \frac{1}{6} \frac{4}{3} = \frac{1}{9} & p_{(12)(13)} &= \frac{1}{2} \frac{1}{6} \frac{4}{3} = \frac{1}{9} & p_{(13)(12)} &= \frac{1}{2} \frac{1}{4} \frac{4}{3} = \frac{1}{6} & p_{(21)(23)} &= \frac{1}{2} \frac{1}{12} \frac{8}{3} = \frac{1}{9} \\
p_{(11)(21)} &= \frac{1}{2} \frac{1}{8} \frac{8}{5} = \frac{1}{10} & p_{(12)(22)} &= \frac{1}{2} \frac{1}{6} \frac{12}{7} = \frac{1}{7} & p_{(13)(23)} &= \frac{1}{2} \frac{1}{12} \frac{3}{1} = \frac{1}{8} & p_{(21)(11)} &= \frac{1}{2} \frac{1}{8} \frac{8}{5} = \frac{4}{15} \\
p_{(11)(31)} &= \frac{1}{2} \frac{1}{6} \frac{8}{5} = \frac{2}{15} & p_{(12)(32)} &= \frac{1}{2} \frac{1}{6} \frac{12}{7} = \frac{1}{7} & p_{(13)(33)} &= \frac{1}{2} \frac{1}{12} \frac{3}{1} = \frac{1}{8} & p_{(21)(31)} &= \frac{1}{2} \frac{1}{6} \frac{8}{5} = \frac{2}{15}
\end{aligned}$$

Edge probabilities.

$$\begin{aligned}
p_{11} p_{(11)(12)} &= \frac{1}{3} \frac{1}{6} = \frac{1}{4} \frac{2}{9} = p_{12} p_{(12)(11)} \\
p_{11} p_{(11)(13)} &= \frac{1}{3} \frac{1}{9} = \frac{1}{6} \frac{2}{9} = p_{13} p_{(13)(11)} \\
p_{11} p_{(11)(21)} &= \frac{1}{3} \frac{1}{10} = \frac{1}{8} \frac{4}{15} = p_{21} p_{(21)(11)}
\end{aligned}$$

Verification of a few edges.

Note that the edge probabilities out of a state such as (1,1) do not add up to one.

That is, with some probability the walk stays at the state that it is in. For example,

$$p_{(11)(11)} = p_{(11)(12)} + p_{(11)(13)} + p_{(11)(21)} + p_{(11)(31)} = 1 - \frac{1}{6} - \frac{1}{24} - \frac{1}{32} - \frac{1}{24} = \frac{9}{32}.$$

Figure 5.10: Using the Gibbs algorithm to set probabilities for a random walk so that the stationary probability will be a desired probability.

from which it follows that $p(\mathbf{x})p_{xy} = p(\mathbf{y})p_{yx}$. By Lemma 5.3 the stationary probability of the random walk is $p(\mathbf{x})$.

5.7 Areas and Volumes

Computing areas and volumes is a classical problem. For many regular figures in two and three dimensions there are closed form formulae. In Chapter 2, we saw how to compute volume of a high dimensional sphere by integration. For general convex sets in d -space, there are no closed form formulae. Can we estimate volumes of d -dimensional convex sets in time that grows as a polynomial function of d ? The MCMC method answers this question in the affirmative.

One way to estimate the area of the region is to enclose it in a rectangle and estimate the ratio of the area of the region to the area of the rectangle by picking random points in the rectangle and seeing what proportion land in the region. Such methods fail in high dimensions. Even for a sphere in high dimension, a cube enclosing the sphere has exponentially larger area, so exponentially many samples are required to estimate the volume of the sphere.

It turns out that the problem of estimating volumes of sets is reducible to the problem of drawing uniform random samples from sets. Suppose one wants to estimate the volume of a convex set R . Create a concentric series of larger and larger spheres $S_1, S_2, \dots, S_k$ such that S_1 is contained in R and S_k contains R . Then

$$\text{Vol}(R) = \text{Vol}(S_k \cap R) = \frac{\text{Vol}(S_k \cap R)}{\text{Vol}(S_{k-1} \cap R)} \frac{\text{Vol}(S_{k-1} \cap R)}{\text{Vol}(S_{k-2} \cap R)} \cdots \frac{\text{Vol}(S_2 \cap R)}{\text{Vol}(S_1 \cap R)} \text{Vol}(S_1)$$

If the radius of the sphere S_i is $1 + \frac{1}{d}$ times the radius of the sphere S_{i-1} , then the value of

$$\frac{\text{Vol}(S_{k-1} \cap R)}{\text{Vol}(S_{k-2} \cap R)}$$

can be estimated by rejection sampling provided one can select points at random from a d -dimensional region. Since the radii of the spheres grows as $1 + \frac{1}{d}$, the number of spheres is at most

$$O(\log_{1+(1/d)} R) = O(Rd).$$

It remains to show how to draw a uniform random sample from a d -dimensional set. It is at this point that we require the set to be convex so that the Markov chain technique we use will converge quickly to its stationary probability. To select a random sample from a d -dimensional convex set impose a grid on the region and do a random walk on the grid points. At each time, pick one of the $2d$ coordinate neighbors of the current grid point, each with probability $1/(2d)$ and go to the neighbor if it is still in the set; otherwise, stay put and repeat. If the grid length in each of the d coordinate directions is at most some a , the total number of grid points in the set is at most a^d . Although this is exponential in

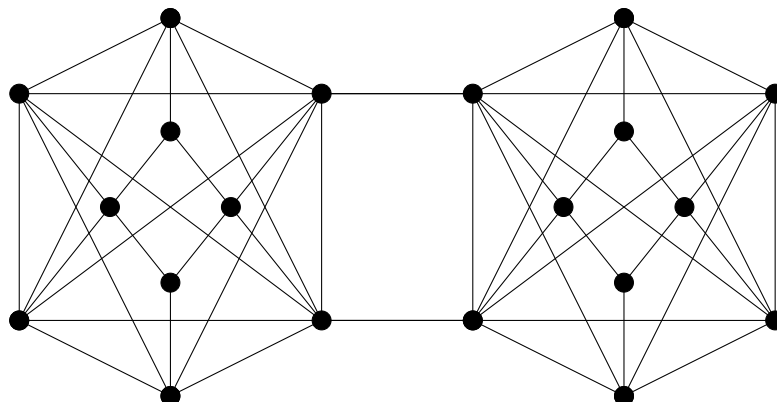


Figure 5.11: A network with a constriction.

d , the Markov chain turns out to be rapidly mixing (the proof is beyond our scope here) and leads to polynomial time bounded algorithm to estimate the volume of any convex set in $\mathbf{R}^d$.

5.8 Convergence of Random Walks on Undirected Graphs

The Metropolis-Hasting algorithm and Gibbs sampling both involve a random walk. Initial states of the walk are highly dependent on the start state of the walk. Both these walks are random walks on edge-weighted undirected graphs. Such Markov chains are derived from electrical networks. Recall the following notation which we will use throughout this section. Given a network of resistors, the conductance of edge (x, y) is denoted c_{xy} and the normalizing constant c_x equals $\sum_y c_{xy}$. The Markov chain has transition probabilities $p_{xy} = c_{xy}/c_x$. We assume the chain is connected. Since

$$c_x p_{xy} = c_x c_{xy} / c_x = c_{xy} = c_{yx} = c_y c_{yx} / c_y = c_y p_{xy}$$

the stationary probabilities are proportional to c_x where the normalization constant is $c_0 = \sum_x c_x$.

An important question is how fast the walk starts to reflect the stationary probability of the Markov process. If the convergence time was proportional to the number of states, the algorithms would not be very useful since the number of states can be exponentially large.

There are clear examples of connected chains that take a long time to converge. A chain with a constriction, see Figure 5.11, takes a long time to converge since the walk is unlikely to reach the narrow passage between the two halves, both of which are reasonably big. We will show in Theorem 5.12 that the time to converge is quantitatively related to

the tightest constriction.

A function is unimodal if it has a single maximum, i.e., it increases and then decreases. A unimodal function like the normal density has no constriction blocking a random walk from getting out of a large set of states, whereas a bimodal function can have a constriction. Interestingly, many common multivariate distributions as well as univariate probability distributions like the normal and exponential are unimodal and sampling according to these distributions can be done using the methods here.

A natural problem is estimating the probability of a convex region in d -space according to a normal distribution. One technique to do this is rejection sampling. Let R be the region defined by the inequality $x_1 + x_2 + \cdots + x_{d/2} \leq x_{d/2+1} + \cdots + x_d$. Pick a sample according to the normal distribution and accept the sample if it satisfies the inequality. If not, reject the sample and retry until one gets a number of samples satisfying the inequality. The probability of the region is approximated by the fraction of the samples that satisfied the inequality. However, suppose R was the region $x_1 + x_2 + \cdots + x_{d-1} \leq x_d$. The probability of this region is exponentially small in d and so rejection sampling runs into the problem that we need to pick exponentially many samples before we accept even one sample. This second situation is typical. Imagine computing the probability of failure of a system. The object of design is to make the system reliable, so the failure probability is likely to be very low and rejection sampling will take a long time to estimate the failure probability.

In general, there could be constrictions that prevent rapid convergence of a Markov chain to its stationary probability. However, if the set is convex in any number of dimensions, then there are no constrictions and there is rapid convergence although the proof of this is beyond the scope of this book.

We define below a combinatorial measure of constriction for a Markov chain, called the *normalized conductance*, and relate this quantity to the rate of convergence to the stationarity probability. One way to avoid constrictions like the one in the picture of Figure 5.11 is to stipulate that the total conductance of edges leaving every subset S of states to $\bar{S}$ be high. The conductance of an edge leaving S is defined to be $\pi_x c_{xy}$ where π_x is the stationary probability of vertex x . This is not possible if S was itself small or even empty. So, in what follows, we “normalize” the total conductance of edges leaving S by the size of S as measured by total c_x for $x \in S$. Recall that $p_{xy} = \frac{c_{xy}}{c_x}$ and the stationary probability $\pi_x = \frac{c_x}{c_0}$ where $c_0 = \sum_x c_x$. In defining the conductance of edges leaving a set we have ignored the normalizing constants.

Definition 5.1 *For a subset S of vertices, the normalized conductance $\Phi(S)$ of S is the*

ratio of the total conductance of all edges from S to $\bar{S}$ to the total of the c_x for $x \in S$.

$$\Phi(S) = \frac{\sum_{\substack{(x,y) \\ x \in S, y \in \bar{S}}} c_{xy}}{\sum_{x \in S} c_x} = \frac{\sum_{\substack{(x,y) \\ x \in S, y \in \bar{S}}} \pi_x p_{xy}}{\sum_{x \in S} \pi_x}$$

■

The normalized conductance⁴ of S is the probability of taking a step from S to outside S conditioned on starting in S in the stationary probability distribution π . The stationary distribution for state x conditioned on being in S is

$$\frac{\pi_x}{\pi(S)} = \frac{c_x}{\sum_{x \in S} c_x}.$$

where $\pi(S) = \sum_{x \in S} \pi_x$.

Definition 5.2 *The normalized conductance of the Markov chain, denoted Φ , is defined by*

$$\Phi = \min_{\substack{S \\ \pi(S) \leq 1/2}} \Phi(S).$$

■

The restriction to sets with $\pi \leq 1/2$ in the definition of Φ is natural. The definition of Φ guarantees that if Φ is high, there is high probability of moving from S to $\bar{S}$ so it is unlikely to get stuck in S provided $\pi(S) \leq \frac{1}{2}$. If $\pi(S) > \frac{1}{2}$, say $\pi(S) = \frac{3}{4}$, then since for every edge $\pi_i p_{ij} = \pi_j p_{ji}$

$$\Phi(S) = \frac{\sum_{i \in S} \pi_i p_{ij}}{\sum_{i \in S} \pi_i} = \frac{\sum_{j \in \bar{S}} \pi_j p_{ji}}{3 \sum_{j \in \bar{S}} \pi_k} = \Phi(\bar{S})/3$$

Since $\Phi(\bar{S}) \geq \Phi$, we still have at least $\Phi/3$ probability of moving out of S . The larger $\pi(S)$ is the smaller the probability of moving out, which is as it should be. We cannot move out of the whole set! One does not need to escape from big sets. Note that a constriction would mean a small Φ .

Definition 5.3 *Fix $\varepsilon > 0$. The ε -mixing time of a Markov chain is the minimum integer t such that for any starting distribution $\mathbf{p}^{(0)}$, the 1-norm distance between the t -step running average probability distribution⁵ and the stationary distribution is at most ε .*

■

⁴We will often drop the word “normalized” and just say “conductance”.

⁵Recall that $\mathbf{a}^{(t)} = \frac{1}{t}(\mathbf{p}^{(0)} + \mathbf{p}^{(1)} + \dots + \mathbf{p}^{(t-1)})$ is called the running average distribution.

The theorem below states that if Φ , the normalized conductance of the Markov chain, is large, then there is fast convergence of the running average probability. Intuitively, if Φ is large, the walk rapidly leaves any subset of states. Later we will see examples where the mixing time is much smaller than the cover time. That is, the number of steps before a random walk reaches a random state independent of its starting state is much smaller than the average number of steps needed to reach every state. In fact for some graphs, called expanders, the mixing time is logarithmic in the number of states.

Theorem 5.12 *The ε -mixing time of a random walk on an undirected graph is*

$$O\left(\frac{\ln(1/\pi_{\min})}{\Phi^2 \varepsilon^3}\right)$$

where $\pi_{\min}$ is the minimum stationary probability of any state.

Proof: Let

$$t = \frac{c \ln(1/\pi_{\min})}{\Phi^2 \varepsilon^2},$$

for a suitable constant c . Let $\mathbf{a} = \mathbf{a}^{(t)}$ be the running average distribution for this value of t . We need to show that $|\mathbf{a} - \boldsymbol{\pi}| \leq \varepsilon$.

Let v_i denote the ratio of the long term average probability for state i at time t divided by the stationary probability for state i . Thus, $v_i = \frac{a_i}{\pi_i}$. Renumber states so that $v_1 \geq v_2 \geq \dots$. A state i for which $v_i > 1$ has more probability than its stationary probability. Execute one step of the Markov chain starting at probabilities $\mathbf{a}$. The probability vector after that step is $\mathbf{a}P$. Now, $\mathbf{a} - \mathbf{a}P$ is the net loss of probability for each state due to the step. Let k be any integer with $v_k > 1$. Let $A = \{1, 2, \dots, k\}$. A is a “heavy” set, consisting of states with $a_i \geq \pi_i$. The net loss of probability for each state from the set A in one step is $\sum_{i=1}^k (a_i - (\mathbf{a}P)_i) \leq \frac{2}{t}$ as in the proof of Theorem 5.2.

Another way to reckon the net loss of probability from A is to take the difference of the probability flow from A to $\bar{A}$ and the flow from $\bar{A}$ to A . For $i < j$,

$$\text{net-flow}(i, j) = \text{flow}(i, j) - \text{flow}(j, i) = \pi_i p_{ij} v_i - \pi_j p_{ji} v_j = \pi_j p_{ji} (v_i - v_j) \geq 0,$$

Thus, for any $l \geq k$, the flow from A to $\{k+1, k+2, \dots, l\}$ minus the flow from $\{k+1, k+2, \dots, l\}$ to A is nonnegative. At each step, heavy sets lose probability. Since for $i \leq k$ and $j > l$, we have $v_i \geq v_k$ and $v_j \leq v_{l+1}$, the net loss from A is at least

$$\sum_{\substack{i \leq k \\ j > l}} \pi_j p_{ji} (v_i - v_j) \geq (v_k - v_{l+1}) \sum_{\substack{i \leq k \\ j > l}} \pi_j p_{ji}.$$

Thus,

$$(v_k - v_{l+1}) \sum_{\substack{i \leq k \\ j > l}} \pi_j p_{ji} \leq \frac{2}{t}.$$

If the total stationary probability $\pi(\{i|v_i \leq 1\})$ of those states where the current probability is less than their stationary probability is less than $\varepsilon/2$, then

$$|\mathbf{a} - \boldsymbol{\pi}|_1 = 2 \sum_{\substack{i \\ v_i \leq 1}} (1 - v_i) \pi_i \leq \varepsilon,$$

so we are done. Assume $\pi(\{i|v_i \leq 1\}) > \varepsilon/2$ so that $\pi(A) \geq \varepsilon \min(\pi(A), \pi(\bar{A}))/2$. Choose l to be the largest integer greater than or equal to k so that

$$\sum_{j=k+1}^l \pi_j \leq \varepsilon \Phi \pi(A)/2.$$

Since

$$\sum_{i=1}^k \sum_{j=k+1}^l \pi_j p_{ji} \leq \sum_{j=k+1}^l \pi_j \leq \varepsilon \Phi \pi(A)/2$$

by the definition of Φ ,

$$\sum_{i \leq k < j} \pi_j p_{ji} \geq \Phi \min(\pi(A), \pi(\bar{A})) \geq \varepsilon \Phi \pi(A).$$

Thus, $\sum_{\substack{i \leq k \\ j > l}} \pi_j p_{ji} \geq \varepsilon \Phi \pi(A)/2$. Substituting into the inequality 5.8 gives

$$v_k - v_{l+1} \leq \frac{8}{t \varepsilon \Phi \pi(A)}. \quad (5.5)$$

This inequality says that v does not drop too much as we go from k to $l+1$. On the other hand, the cumulative total of π will have increased, since, $\pi_1 + \pi_2 + \dots + \pi_{l+1} \geq \rho(\pi_1 + \pi_2 + \dots + \pi_k)$, where, $\rho = 1 + \frac{\varepsilon \Phi}{2}$. We will be able to use this repeatedly to argue that overall v does not drop too much. If that is the case (in the extreme, for example, if all the v_i are 1 each), then intuitively, $\mathbf{a} \approx \boldsymbol{\pi}$, which is what we are trying to prove. Unfortunately, the technical execution of this argument is a bit messy - we have to divide $\{1, 2, \dots, n\}$ into groups and consider the drop in v as we move from one group to the next and then add up. We do this now.

Now, divide $\{1, 2, \dots\}$ into groups as follows. The first group G_1 is $\{1\}$. In general, if the r^{th} group G_r begins with state k , the next group G_{r+1} begins with state $l+1$ where l is as defined above. Let i_0 be the largest integer with $v_{i_0} > 1$. Stop with G_m , if G_{m+1} would begin with an $i > i_0$. If group G_r begins in i , define $u_r = v_i$.

$$|\mathbf{a} - \boldsymbol{\pi}|_1 \leq 2 \sum_{i=1}^{i_0} \pi_i (v_i - 1) \leq \sum_{r=1}^m \pi(G_r) (u_r - 1) = \sum_{r=1}^m \pi(G_1 \cup G_2 \cup \dots \cup G_r) (u_r - u_{r+1}),$$

where the analog of integration by parts for sums is used in the last step using the convention that $u_{m+1} = 1$. Since $u_r - u_{r+1} \leq 8/\varepsilon \Phi \pi(G_1 \cup \dots \cup G_r)$, the sum is at most $8m/t\varepsilon \Phi$. Since $\pi_1 + \pi_2 + \dots + \pi_{l+1} \geq \rho(\pi_1 + \pi_2 + \dots + \pi_k)$,

$$m \leq \ln_\rho(1/\pi_1) \leq \ln(1/\pi_1)/(\rho - 1).$$

Thus $|\mathbf{a} - \boldsymbol{\pi}|_1 \leq O(\ln(1/\pi_{\min})/t\Phi^2\varepsilon^2) \leq \varepsilon$ for a suitable choice of c and this completes the proof. $\blacksquare$

5.8.1 Using Normalized Conductance to Prove Convergence

We now apply Theorem 5.12 to some examples to illustrate how the normalized conductance bounds the rate of convergence. Our first examples will be simple graphs. The graphs do not have rapid converge, but their simplicity helps illustrate how to bound the normalized conductance and hence the rate of convergence.

A 1-dimensional lattice

Consider a random walk on an undirected graph consisting of an n -vertex path with self-loops at the both ends. With the self loops, the stationary probability is a uniform $\frac{1}{n}$ over all vertices. The set with minimum normalized conductance is the set with the maximum number of vertices with the minimum number of edges leaving it. This set consists of the first $n/2$ vertices, for which total conductance of edges from S to $\bar{S}$ is $\pi_{n/2} p_{n/2, n/2+1} = \Omega(\frac{1}{n})$ and $\pi(S) = \frac{1}{2}$. ($\pi_{n/2}$ is the stationary probability of vertex numbered $\frac{n}{2}$.) Thus

$$\Phi(S) = 2\pi_{n/2} p_{n/2, n/2+1} = \Omega(1/n).$$

By Theorem 5.12, for ε a constant such as $1/100$, after $O(n^2 \log n)$ steps, $|\mathbf{a}^{(t)} - \boldsymbol{\pi}|_1 \leq 1/100$. This graph does not have rapid convergence. The hitting time and the cover time are $O(n^2)$. In many interesting cases, the mixing time may be much smaller than the cover time. We will see such an example later.

A 2-dimensional lattice

Consider the $n \times n$ lattice in the plane where from each point there is a transition to each of the coordinate neighbors with probability $1/4$. At the boundary there are self-loops with probability $1 - (\text{number of neighbors})/4$. It is easy to see that the chain is connected. Since $p_{ij} = p_{ji}$, the function $f_i = 1/n^2$ satisfies $f_i p_{ij} = f_j p_{ji}$ and by Lemma 5.3 is the stationary probability. Consider any subset S consisting of at most half the states. Index states by their x and y coordinates. For at least half the states in S , either row x or column y intersects $\bar{S}$ (Exercise 5.46). So at least $\Omega(|S|/n)$ points in S are adjacent to points in $\bar{S}$. Each such point contributes $\pi_i p_{ij} = \Omega(1/n^2)$ to $\text{flow}(S, \bar{S})$. So

$$\sum_{i \in S} \sum_{j \in \bar{S}} \pi_i p_{ij} \geq c|S|/n^3.$$

Thus, $\Phi \geq \Omega(1/n)$. By Theorem 5.12, after $O(n^2 \ln n / \varepsilon^2)$ steps, $|\mathbf{a}^{(t)} - \boldsymbol{\pi}|_1 \leq 1/100$.

A lattice in d -dimensions

Next consider the $n \times n \times \cdots \times n$ lattice in d -dimensions with a self-loop at each boundary point with probability $1 - (\text{number of neighbors})/2d$. The self loops make all π_i equal to n^{-d} . View the lattice as an undirected graph and consider the random walk on this undirected graph. Since there are n^d states, the cover time is at least n^d and thus exponentially dependent on d . It is possible to show (Exercise 5.62) that Φ is $\Omega(1/dn)$. Since all π_i are equal to n^{-d} , the mixing time is $O(d^3 n^2 \ln n / \varepsilon^2)$, which is polynomially bounded in n and d .

The d -dimensional lattice is related to the Metropolis-Hastings algorithm and Gibbs sampling although in those constructions there is a nonuniform probability distribution at the vertices. However, the d -dimension lattice case suggests why the Metropolis-Hastings and Gibbs sampling constructions might converge fast.

A clique

Consider an n vertex clique. The stationary probability of each vertex is $\frac{1}{n}$. Thus

$$\sum_{x \in S} c_x = \frac{|S|}{n}.$$

The probability of each edge is $\frac{1}{n}$ if we have self loops at each vertex. Thus

$$\sum_{(x,y)} \pi_x c_{xy} = \frac{1}{n} \frac{1}{n} |S| |\bar{S}|$$

and $\Phi(S) = \frac{|\bar{S}|}{n}$. $\Phi = \min \Phi(S)$ for $|S| \leq \frac{n}{2}$ and hence $|\bar{S}| \geq \frac{n}{2}$. Thus $\Phi = \frac{1}{2}$. The mixing time is

$$O\left(\frac{\ln \frac{1}{\pi_{\min}}}{\Phi^2 \epsilon^3}\right) = O\left(\frac{\ln n}{\frac{1}{4} \epsilon^3}\right) = O(\ln n).$$

A connected undirected graph

Next consider a random walk on a connected n vertex undirected graph where at each vertex all edges are equally likely. The stationary probability of a vertex equals the degree of the vertex divided by the sum of degrees which equals twice the number of edges. The sum of the vertex degrees is at most n^2 and thus, the steady state probability of each vertex is at least $\frac{1}{n^2}$. Since the degree of a vertex is at most n , the probability of each edge at a vertex is at least $\frac{1}{n}$. For any S , the total conductance of edges out of S is greater than or equal to

$$\frac{1}{n^2} \frac{1}{n} = \frac{1}{n^3}.$$

Thus, Φ is at least $\frac{1}{n^3}$. Since $\pi_{\min} \geq \frac{1}{n^2}$, $\ln \frac{1}{\pi_{\min}} = O(\ln n)$. Thus, the mixing time is $O(n^6(\ln n)/\varepsilon^2)$.

The Gaussian distribution on the interval $[-1, 1]$

Consider the interval $[-1, 1]$. Let δ be a “grid size” specified later and let G be the graph consisting of a path on the $\frac{2}{\delta} + 1$ vertices $\{-1, -1 + \delta, -1 + 2\delta, \dots, 1 - \delta, 1\}$ having self loops at the two ends. Let $\pi_x = ce^{-\alpha x^2}$ for $x \in \{-1, -1 + \delta, -1 + 2\delta, \dots, 1 - \delta, 1\}$ where $\alpha > 1$ and c has been adjusted so that $\sum_x \pi_x = 1$.

We now describe a simple Markov chain with the π_x as its stationary probability and argue its fast convergence. With the Metropolis-Hastings’ construction, the transition probabilities are

$$p_{x, x+\delta} = \frac{1}{2} \min \left(1, \frac{e^{-\alpha(x+\delta)^2}}{e^{-\alpha x^2}} \right) \text{ and } p_{x, x-\delta} = \frac{1}{2} \min \left(1, \frac{e^{-\alpha(x-\delta)^2}}{e^{-\alpha x^2}} \right).$$

Let S be any subset of states with $\pi(S) \leq \frac{1}{2}$. Consider the case when S is an interval $[k\delta, 1]$ for $k \geq 1$. It is easy to see that

$$\begin{aligned} \pi(S) &\leq \int_{x=(k-1)\delta}^{\infty} ce^{-\alpha x^2} dx \\ &\leq \int_{(k-1)\delta}^{\infty} \frac{x}{(k-1)\delta} ce^{-\alpha x^2} dx \\ &= O \left(\frac{ce^{-\alpha((k-1)\delta)^2}}{\alpha(k-1)\delta} \right). \end{aligned}$$

Now there is only one edge from S to $\bar{S}$ and total conductance of edges out of S is

$$\sum_{i \in S} \sum_{j \notin S} \pi_i p_{ij} = \pi_{k\delta} p_{k\delta, (k-1)\delta} = \min(ce^{-\alpha k^2 \delta^2}, ce^{-\alpha(k-1)^2 \delta^2}) = ce^{-\alpha k^2 \delta^2}.$$

Using $1 \leq k \leq 1/\delta$ and $\alpha \geq 1$, $\Phi(S)$ is

$$\begin{aligned} \Phi(S) &= \frac{\text{flow}(S, \bar{S})}{\pi(S)} \geq ce^{-\alpha k^2 \delta^2} \frac{\alpha(k-1)\delta}{ce^{-\alpha((k-1)\delta)^2}} \\ &\geq \Omega(\alpha(k-1)\delta e^{-\alpha \delta^2(2k-1)}) \geq \Omega(\delta e^{-O(\alpha \delta)}). \end{aligned}$$

For $\delta < \frac{1}{\alpha}$, we have $\alpha \delta < 1$, so $e^{-O(\alpha \delta)} = \Omega(1)$, thus, $\Phi(S) \geq \Omega(\delta)$. Now, $\pi_{\min} \geq ce^{-\alpha} \geq e^{-1/\delta}$, so $\ln(1/\pi_{\min}) \leq 1/\delta$.

If S is not an interval of the form $[k, 1]$ or $[-1, k]$, then the situation is only better since there is more than one “boundary” point which contributes to $\text{flow}(S, \bar{S})$. We do

not present this argument here. By Theorem 5.12 in $\Omega(1/\delta^3\varepsilon^2)$ steps, a walk gets within ε of the steady state distribution.

In these examples, we have chosen simple probability distributions. The methods extend to more complex situations.

5.9 Bibliographic Notes

The material on the analogy between random walks on undirected graphs and electrical networks is from [DS84] as is the material on random walks in Euclidean space. Additional material on Markov chains can be found in [MR95b], [MU05], and [per10]. For material on Markov Chain Monte Carlo methods see [Jer98] and [Liu01].

The use of normalized conductance to prove convergence of Markov Chains is by Sinclair and Jerrum, [SJ] and Alon [Alo86]. A polynomial time bounded Markov chain based method for estimating the volume of convex sets was developed by Dyer, Frieze and Kannan [DFK91].

5.10 Exercises

Exercise 5.1 *The Fundamental Theorem of Markov chains proves that for a connected Markov chain, the long-term average distribution $\mathbf{a}^{(t)}$ converges to a stationary distribution. Does the t step distribution $\mathbf{p}^{(t)}$ also converge for every connected Markov Chain? Consider the following examples : (i) A two-state chain with $p_{12} = p_{21} = 1$. (ii) A three state chain with $p_{12} = p_{23} = p_{31} = 1$ and the other $p_{ij} = 0$. Generalize these examples to produce Markov Chains with many states.*

A connected Markov Chain is said to be aperiodic if the greatest common divisor of the lengths of directed cycles is 1. It is known (though we do not prove it here) that for connected aperiodic chains, $\mathbf{p}^{(t)}$ converges to the stationary distribution.

Exercise 5.2

1. *What is the set of possible harmonic functions on a connected graph if there are only interior vertices and no boundary vertices that supply the boundary condition?*
2. *Let q_x be the stationary probability of vertex x in a random walk on an undirected graph where all edges at a vertex are equally likely and let d_x be the degree of vertex x . Show that $\frac{q_x}{d_x}$ is a harmonic function.*
3. *If there are multiple harmonic functions when there are no boundary conditions why is the stationary probability of a random walk on an undirected graph unique?*
4. *What is the stationary probability of a random walk on an undirected graph?*

Exercise 5.3 *In Section ?? we associate a graph and edge probabilities with an electric network such that voltages and currents in the electrical network corresponded to properties of random walks on the graph. Can we go in the reverse order and construct the equivalent electrical network from a graph with edge probabilities?*

Probabilistic interpretation of voltages

Exercise 5.4 *Given a graph consisting of a single path of five vertices numbered 1 to 5, what is the probability of reaching vertex 1 before vertex 5 when starting at vertex 4.*

Exercise 5.5 *Consider the electrical resistive network in Figure 5.12 consisting of vertices connected by resistors. Kirchoff's law states that the currents at each vertex sum to zero. Ohm's law states that the voltage across a resistor equals the product of the resistance times the current through it. Using these laws calculate the effective resistance of the network.*

Probabilistic interpretation of current

Exercise 5.6 *Consider the electrical network of Figure 5.13.*

1. *Set the voltage at a to one and at b to zero. What are the voltages at c and d ?*

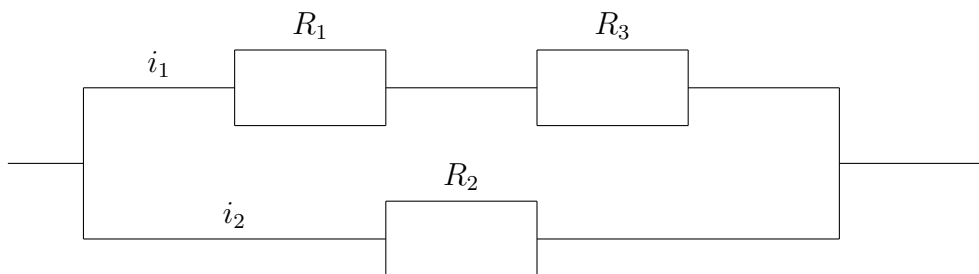


Figure 5.12: An electrical network of resistors.

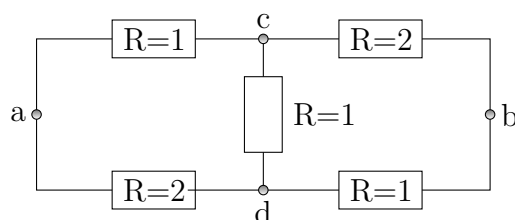


Figure 5.13: An electrical network of resistors.

2. What is the current in the edges a to c , a to d , c to d , c to b and d to b ?
3. What is the effective resistance between a and b ?
4. Convert the electrical network to a graph. What are the edge probabilities at each vertex?
5. What is the probability of a walk starting at c reaching a before b ? a walk starting at d ?
6. What is the net frequency that a walk from a to b goes through the edge from c to d ?
7. What is the probability that a random walk starting at a will return to a before reaching b ?

Effective Resistance and Escape Probability

Exercise 5.7 Consider a graph corresponding to an electrical network with vertices a and b . Prove directly that $\frac{c_{eff}}{c_a}$ must be less than or equal to one. We know that this is the escape probability and must be at most 1. But, for this exercise, do not use that fact.

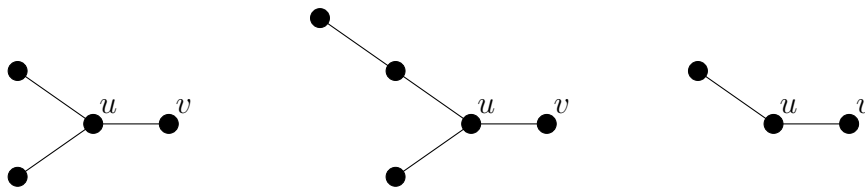


Figure 5.14: Three graphs

Exercise 5.8 (*Thomson's Principle*) The energy dissipated by the resistance of edge xy in an electrical network is given by $i_{xy}^2 r_{xy}$. The total energy dissipation in the network is $E = \frac{1}{2} \sum_{x,y} i_{xy}^2 r_{xy}$ where the $\frac{1}{2}$ accounts for the fact that the dissipation in each edge is counted twice in the summation. Show that the actual current distribution is that distribution satisfying Ohm's law that minimizes energy dissipation.

Exercise 5.9 (*Rayleigh's law*) Prove that reducing the value of a resistor in a network cannot increase the effective resistance. Prove that increasing the value of a resistor cannot decrease the effective resistance. You may use Thomson's principle 5.8.

Random walks on undirected graphs hitting time

Exercise 5.10 What is the hitting time h_{uv} for two adjacent vertices on a cycle of length n ? What is the hitting time if the edge (u, v) is removed?

Exercise 5.11 What is the hitting time h_{uv} for the three graphs in Figure 5.14.

Exercise 5.12 Show that adding an edge can either increase or decrease hitting time by calculating h_{24} for the three graphs in Figure 5.15.

Exercise 5.13 Consider the n vertex connected graph shown in Figure 5.16 consisting of an edge (u, v) plus a connected graph on $n - 1$ vertices and some number of edges. Prove that $h_{uv} = 2m - 1$ where m is the number of edges in the $n - 1$ vertex subgraph.

Exercise 5.14 What is the most general solution to the difference equation $t(i + 2) - 5t(i + 1) + 6t(i) = 0$. How many boundary conditions do you need to make the solution unique?

Exercise 5.15 Given the difference equation $a_k t(i + k) + a_{k-1} t(i + k - 1) + \cdots + a_1 t(i + 1) + a_0 t(i) = 0$ the polynomial $a_k t^k + a_{k-1} t^{k-1} + \cdots + a_1 t + a_0 = 0$ is called the characteristic polynomial.

1. If the equation has a set of r distinct roots, what is the most general form of the solution?

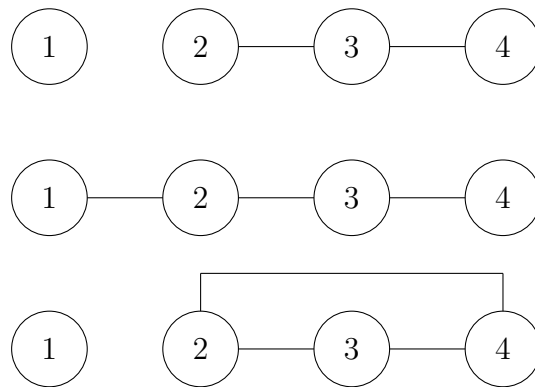


Figure 5.15: Three graph

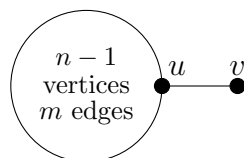


Figure 5.16: A connected graph consisting of $n - 1$ vertices and m edges along with a single edge (u, v) .

2. If the roots of the characteristic polynomial are not unique what is the most general form of the solution?
3. What is the dimension of the solution space?
4. If the difference equation is not homogeneous and $f(i)$ is a specific solution to the nonhomogeneous difference equation, what is the full set of solutions to the difference equation?

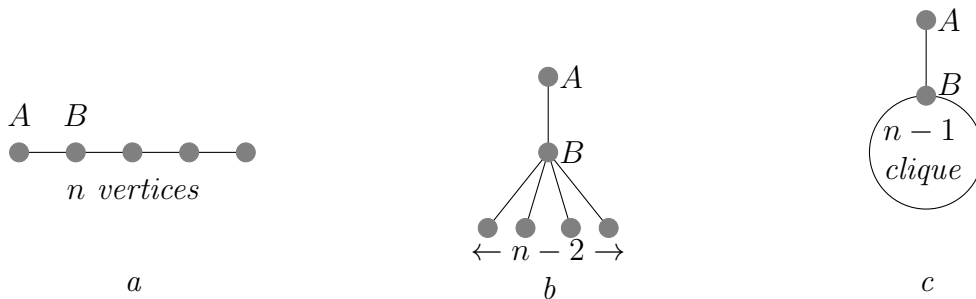
Exercise 5.16 Given the integers 1 to n , what is the expected number of draws with replacement until the integer 1 is drawn.

Exercise 5.17 Consider the set of integers $\{1, 2, \dots, n\}$. What is the expected number of draws d with replacement so that every integer is drawn?

Exercise 5.18 Consider a random walk on a clique of size n . What is the expected number of steps before a given vertex is reached?

Exercise 5.19 Show that adding an edge to a graph can either increase or decrease commute time.

Exercise 5.20 For each of the three graphs below what is the return time starting at vertex A ? Express your answer as a function of the number of vertices, n , and then express it as a function of the number of edges m .



Exercise 5.21 Suppose that the clique in Exercise 5.20 was an arbitrary graph with $m-1$ edges. What would be the return time to A in terms of m , the total number of edges.

Exercise 5.22 Suppose that the clique in Exercise 5.20 was an arbitrary graph with $m-d$ edges and there were d edges from A to the graph. What would be the expected length of a random path starting at A and ending at A after returning to A exactly d times.

Exercise 5.23 Given an undirected graph with a component consisting of a single edge find two eigenvalues of the Laplacian $L = D - A$ where D is a diagonal matrix with vertex degrees on the diagonal and A is the adjacency matrix of the graph.

Exercise 5.24 A researcher was interested in determining the importance of various edges in an undirected graph. He computed the stationary probability for a random walk on the graph and let p_i be the probability of being at vertex i . If vertex i was of degree d_i , the frequency that edge (i, j) was traversed from i to j would be $\frac{1}{d_i}p_i$ and the frequency that the edge was traversed in the opposite direction would be $\frac{1}{d_j}p_j$. Thus, he assigned an importance of $\left| \frac{1}{d_i}p_i - \frac{1}{d_j}p_j \right|$ to the edge. What is wrong with his idea?

Exercise 5.25 Prove that two independent random walks on a two dimensional lattice will meet with probability one.

Exercise 5.26 Suppose two individuals are flipping balanced coins and each is keeping track of the number of heads minus the number of tails. Will both individual's count return to zero at the same time?

Exercise 5.27 Consider the lattice in 2-dimensions. In each square add the two diagonal edges. What is the escape probability for the resulting graph?

Exercise 5.28 Determine by simulation the escape probability for the 3-dimensional lattice.

Exercise 5.29 What is the escape probability for a random walk starting at the root of a binary tree?

Exercise 5.30 Consider a random walk on the positive half line, that is the integers $1, 2, 3, \dots$. At the origin, always move right one step. At all other integers move right with probability $2/3$ and left with probability $1/3$. What is the escape probability?

Exercise 5.31 ** What is the probability of returning to the start vertex on a random walk on an infinite planar graph?

Exercise 5.32 Create a model for a graph similar to a 3-dimensional lattice in the way that a planar graph is similar to a 2-dimensional lattice. What is probability of returning to the start vertex in your model?

Random walks on directed graphs

Exercise 5.33 Consider the graphs in Figure 5.17. Calculate the stationary distribution for a random walk on each graph and the flow through each edge. What condition holds on the flow through edges in the undirected graph? In the directed graph?

Finite markov processes Page rank and hitting time

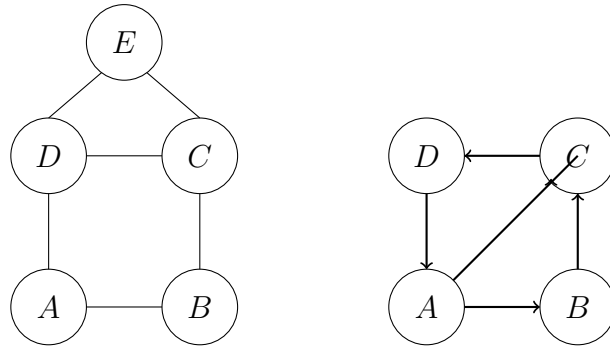


Figure 5.17: An undirected and a directed graph.

Exercise 5.34 Create a random directed graph with 200 vertices and roughly eight edges per vertex. Add k new vertices and calculate the page rank with and without directed edges from the k added vertices to vertex 1. How much does adding the k edges change the page rank of vertices for various values of k and restart frequency? How much does adding a loop at vertex 1 change the page rank? To do the experiment carefully one needs to consider the page rank of a vertex to which the star is attached. If it has low page rank its page rank is likely to increase a lot.

Exercise 5.35 Repeat the experiment in Exercise 5.34 for hitting time.

Exercise 5.36 Search engines ignore self loops in calculating page rank. Thus, to increase page rank one needs to resort to loops of length two. By how much can you increase the page rank of a page by adding a number of loops of length two?

Exercise 5.37 Can one increase the page rank of a vertex v in a directed graph by doing something some distance from v ? The answer is yes if there is a long narrow chain of vertices into v with no edges leaving the chain. What if there is no such chain?

Exercise 5.38 Given a very large directed graph with many vertices of out degree one or in degree one, can one compute page rank of a reduced graph in which the vertices of in degree or out degree one have been merged and then compute the page rank of the original graph from the pagerank of the reduced graph? Does the method work if there are random restarts?

Exercise 5.39 If we model random restarts by adding a restart vertex, do we get the same results as if we eliminated the restart vertex and added n^2 edges?

Exercise 5.40 Consider modifying personal page rank as follows. Start with the uniform restart distribution and calculate the steady state probabilities. Then run the personalized

page rank algorithm using the stationary distribution calculated instead of the uniform distribution. Keep repeating until the process converges. That is, we get a stationary probability distribution such that if we use the stationary probability distribution for the restart distribution we will get the stationary probability distribution back. Does this process converge? What is the resulting distribution? What distribution do we get for the graph consisting of two vertices u and v with a single edge from u to v ?

Exercise 5.41 Number the vertices of a graph $\{1, 2, \dots, n\}$. Define hitting time to be the expected time from vertex 1. In (2) assume that the vertices in the cycle are sequentially numbered.

1. What is the hitting time for a vertex in a complete directed graph with self loops?
2. What is the hitting time for a vertex in a directed cycle with n vertices?

Create exercise relating strongly connected and full rank

Full rank implies strongly connected.

Strongly connected does not necessarily imply full rank

$$\begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 1 \\ 1 & 1 & 0 \end{pmatrix}$$

Is graph aperiodic iff $\lambda_1 > \lambda_2$?

Exercise 5.42 Using a web browser bring up a web page and look at the source html. How would you extract the url's of all hyperlinks on the page if you were doing a crawl of the web? With Internet Explorer click on "source" under "view" to access the html representation of the web page. With Firefox click on "page source" under "view".

Exercise 5.43 Sketch an algorithm to crawl the World Wide Web. There is a time delay between the time you seek a page and the time you get it. Thus, you cannot wait until the page arrives before starting another fetch. There are conventions that must be obeyed if one were to actually do a search. Sites specify information as to how long or which files can be searched. Do not attempt an actual search without guidance from a knowledgeable person.

Exercise 5.44 Let $p(\mathbf{x})$, where $\mathbf{x} = (x_1, x_2, \dots, x_d)$ $x_i \in \{0, 1\}$, be a multivariate probability distribution. For $d = 100$, how would you estimate the marginal distribution

$$p(\mathbf{x}) = \sum_{x_2, \dots, x_d} p(x_1, x_2, \dots, x_d)?$$

Markov Chain Monte Carlo Time reversibility

Exercise 5.45 Prove Proposition 5.11 that for two probability distributions $\mathbf{p}, \mathbf{q}$, $\|\mathbf{p} - \mathbf{q}\|_1 = 2 \sum_i (p_i - q_i)^+$.

Exercise 5.46 Suppose S is a subset of at most $n^2/2$ points in the $n \times n$ lattice. Show that

$$|\{(i, j) \in S \mid \text{all elements in row } i \text{ and all elements in column } j \text{ are in } S\}| \leq |S|/2.$$

Exercise 5.47 Show that the stationary probabilities of the chain described in the Gibbs sampler is the correct p .

Exercise 5.48 A Markov chain is said to be symmetric if for all i and j , $p_{ij} = p_{ji}$. What is the stationary distribution of a connected symmetric chain? Prove your answer.

Exercise 5.49 How would you integrate a multivariate polynomial distribution over some region?

Exercise 5.50 Given a time-reversible Markov chain, modify the chain as follows. At the current state, stay put (no move) with probability $1/2$. With the other probability $1/2$, move as in the old chain. Show that the new chain has the same stationary distribution. What happens to the convergence time in this modification?

Metropolis Hasting algorithm

Exercise 5.51 Using the Metropolis-Hasting Algorithm create a Markov chain whose stationary probability is that given in the following table.

$x_1 x_2$	00	01	02	10	11	12	20	21	22
Prob	1/16	1/8	1/16	1/8	1/4	1/8	1/16	1/8	1/16

Exercise 5.52 Let $\mathbf{p}$ be a probability vector (nonnegative components adding up to 1) on the vertices of a connected graph. Set p_{ij} (the transition probability from i to j) to p_j for all $i \neq j$ which are adjacent in the graph. Show that the stationary probability vector for the chain is $\mathbf{p}$. Is running this chain an efficient way to sample according to a distribution close to $\mathbf{p}$? Think, for example, of the graph G being the $n \times n \times n \times \cdots \times n$ grid.

Exercise 5.53 Construct the edge probability for a three state Markov chain where each pair of states is connected by an edge so that the stationary probability is $(\frac{1}{2}, \frac{1}{3}, \frac{1}{6})$.

Exercise 5.54 Consider a three state Markov chain with stationary probability $(\frac{1}{2}, \frac{1}{3}, \frac{1}{6})$. Consider the Metropolis-Hastings algorithm with G the complete graph on these three vertices. What is the expected probability that we would actually make a move along a selected edge?

Gibbs sampling

Exercise 5.55 Try Gibbs sampling on $p(x) = \begin{pmatrix} \frac{1}{2} & 0 \\ 0 & \frac{1}{2} \end{pmatrix}$.

What happens? How does the Metropolis Hasting Algorithm do?

Exercise 5.56 Consider $p(\mathbf{x})$, where, $\mathbf{x} = (x_1, \dots, x_{100})$ and $p(\mathbf{0}) = \frac{1}{2}$, $p(\mathbf{x}) = \frac{1}{(2^{100}-1)}$ $x \neq \mathbf{0}$. How does Gibbs sampling behave?

Convergence to steady state

Exercise 5.57 Construct an algorithm and compute the volume of a unit radius sphere in 20 dimensions by carrying out a random walk on a 20 dimensional grid with 0.1 spacing.

Exercise 5.58 Given a graph G and an integer k how would you generate connected subgraphs of G with k vertices with probability proportional to the number of edges in the subgraph induced on those vertices? The probabilities need not be exactly proportional to the number of edges and you are not expected to prove your algorithm for this problem.

Exercise 5.59 Suppose one wishes to generate uniformly at random regular, degree three undirected, connected multi-graphs each with 1,000 vertices. A multi-graph may have multiple edges between a pair of vertices and self loops. One decides to do this by a Markov Chain Monte Carlo technique. They design a network where each vertex is a regular degree three, 1,000 vertex multi-graph. For edges they say that the vertices corresponding to two graphs are connected by an edge if one graph can be obtained from the other by a flip of a pair of disjoint edges. In a flip, a pair of edges (a, b) and (c, d) are replaced by (a, c) and (b, d) .

1. Prove that a swap on a connected multi-graph results in a connected multi-graph.
2. Prove that the network whose vertices correspond to the desired graphs is connected.
3. Prove that the stationary probability of the random walk is uniform.
4. Give an upper bound on the diameter of the network.

In order to use a random walk to generate the graphs uniformly at random, the random walk must rapidly converge to the stationary probability. Proving this is beyond the material in this book.

Exercise 5.60 What is the mixing time for

1. Two cliques connected by a single edge?
2. A graph consisting of an n vertex clique plus one additional vertex connected to one vertex in the clique.

Exercise 5.61 What is the mixing time for

1. $G(n, p)$ with $p = \frac{\log n}{n}$?
2. a circle with n vertices where at each vertex an edge has been added to another vertex chosen at random. On average each vertex will have degree four, two circle edges, and edge from that vertex to a vertex chosen at random, and possibly some edges that are the ends of the random edges from other vertices.

Using minimum escape probability to prove convergence

Exercise 5.62 Show that for the $n \times n \times \cdots \times n$ grid in d space, the normalized conductance is $\Omega(1/dn)$.

Hint: The argument is a generalization of the argument in Exercise 5.46. Argue that for any subset S containing at most $1/2$ the grid points, for at least $1/2$ the grid points in S , among the d coordinate lines through the point, at least one intersects $\bar{S}$.

6 Learning and VC-dimension

6.1 Learning

Learning algorithms are general purpose tools that solve problems from many domains without detailed domain-specific knowledge. They have proven to be very effective in a large number of contexts. The task of a learning algorithm is to learn to classify a set of objects. To illustrate with an example, suppose one wants an algorithm to distinguish among different types of motor vehicles such as cars, trucks, and tractors. Using domain knowledge about motor vehicles, one can create a set of *features*. Some examples of features are the number of wheels, the power of the engine, the number of doors, and the length of vehicle. If there are d features, each object can be represented as a d -dimensional vector, called the feature vector, with each component of the vector giving the value of one feature. The objective is to design a “prediction” algorithm that given a vector will correctly predict the corresponding type of vehicle. Earlier rule-based approaches to this problem used domain knowledge to develop a set of rules such as: if the number of wheels is four, it is a car. Prediction was done by checking the rules.

In the learning approach, the process of developing the prediction rules is not domain-specific; it is automated. In learning, domain expertise is used to decide on the choice of features, reducing the problem to one of classifying feature vectors. Further, a domain expert is called on to classify a set of feature vectors, called *training examples*, and present these as input to the learning algorithm. The role of the expert ends here.

The learning algorithm takes as input the set of labeled training examples and develops a set of rules that applied to the training vectors gives the correct labels. In the motor vehicle example, the learning algorithm needs no knowledge of this domain at all. It just deals with a set of training vectors in d -dimensional space and produces a rule to classify d -dimensional space into regions, one region corresponding to each of “car”, “truck”, etc.

The task of the learning algorithm is to output a set of rules that correctly labels all training examples. Of course, for this limited task, one could output the rule “for each training example, use the label that the expert has already supplied”. But, we insist on *Occam’s razor principle* that states that the rules output by the algorithm, must be more succinct than the table of all labeled training examples. This is akin to developing a scientific theory to explain extensive observations. The theory must be more succinct than just a list of observations.

The general task is not to be correct just on the training examples, but have the learnt rules correctly predict the labels of future examples. Intuitively, if the classifier is trained on sufficiently many training examples, then it seems likely that it would work well on the space of all examples. We will see later that the theory of Vapnik-Chervonenkis dimension (VC-dimension) confirms this intuition. For now, our attention is focussed on getting a

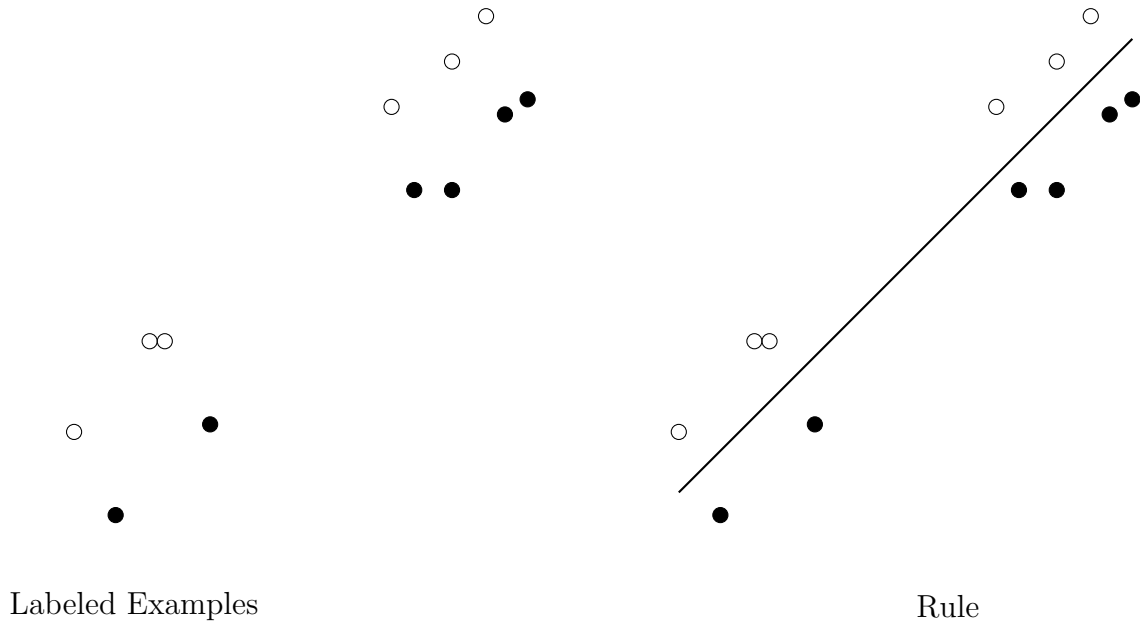


Figure 6.1: Training set and the rule that is learnt

succinct set of rules that correctly classifies the training examples. This is referred to as “learning”.

Throughout this chapter, we assume all the labels are binary. It is not difficult to see that the general problem of classifying into one of several types can be reduced to binary classification. Classifying into car or non-car, tractor or non-tractor, etc. will pin down the type of vehicle. So the teacher’s labels are assumed to be $+1$ or -1 . For an illustration, see Figure 6.1 where examples are in 2-dimensions corresponding to two features. Examples labeled -1 are unfilled circles and those labeled $+1$ are filled circles. The right hand picture illustrates a rule that the algorithm could come up with, the examples above the line are -1 and those below are $+1$.

The simplest rule in d -dimensional space is the generalization of a line in the plane, namely, a half-space. Does a weighted sum of feature values exceed a threshold? Such a rule may be thought of as being implemented by a threshold gate that takes the feature values as inputs, computes their weighted sum and outputs yes or no depending on whether or not the sum is greater than the threshold. One could also look at a network of interconnected threshold gates called a neural net. Threshold gates are sometimes called perceptrons since one model of human perception is that it is done by a neural net in the brain.

6.2 Linear Separators, the Perceptron Algorithm, and Margins

The problem of learning a half-space or a linear separator consists of n labeled examples, $\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n$, in d -dimensional space. The task is to find a d -dimensional vector $\mathbf{w}$, if one exists, and a threshold b such that

$$\begin{aligned} \mathbf{w}^T \mathbf{a}_i &> b \text{ for each } \mathbf{a}_i \text{ labelled } +1 \\ \mathbf{w}^T \mathbf{a}_i &< b \text{ for each } \mathbf{a}_i \text{ labelled } -1. \end{aligned} \tag{6.1}$$

A vector-threshold pair, $(\mathbf{w}, b)$, satisfying the inequalities is called a linear separator.

The above formulation is a linear program (LP) in the unknowns $\mathbf{w}$ and b that can be solved by a general purpose LP algorithm. Linear programming is solvable in polynomial time but a simpler algorithm called the perceptron learning algorithm can be much faster when there is a feasible solution $\mathbf{w}$ with a lot of wiggle room or margin, though it is not polynomial time bounded in general.

We begin by adding an extra coordinate to each $\mathbf{a}_i$ and $\mathbf{w}$, writing $\hat{\mathbf{a}}_i = (\mathbf{a}_i, 1)$ and $\hat{\mathbf{w}} = (\mathbf{w}, -b)$. Suppose l_i is the ± 1 label on $\mathbf{a}_i$. Then, the inequalities in (6.1) can be rewritten as

$$(\hat{\mathbf{w}}^T \hat{\mathbf{a}}_i) l_i > 0 \quad 1 \leq i \leq n.$$

Since the right hand side is zero, we may scale $\hat{\mathbf{a}}_i$ so that $|\hat{\mathbf{a}}_i| = 1$. Adding the extra coordinate increased the dimension by one but now the separator contains the origin. For simplicity of notation, in the rest of this section, we drop the hats and let $\mathbf{a}_i$ and $\mathbf{w}$ stand for the corresponding $\hat{\mathbf{a}}_i$ and $\hat{\mathbf{w}}$.

The perceptron learning algorithm

The perceptron learning algorithm is simple and elegant. We wish to find a solution $\mathbf{w}$ to

$$(\mathbf{w}^T \mathbf{a}_i) l_i > 0 \quad 1 \leq i \leq n \tag{6.2}$$

where $|\mathbf{a}_i| = 1$. Starting with $\mathbf{w} = l_1 \mathbf{a}_1$, pick any $\mathbf{a}_i$ with $(\mathbf{w}^T \mathbf{a}_i) l_i \leq 0$, and replace $\mathbf{w}$ by $\mathbf{w} + l_i \mathbf{a}_i$. Repeat until $(\mathbf{w}^T \mathbf{a}_i) l_i > 0$ for all i .

The intuition behind the algorithm is that correcting $\mathbf{w}$ by adding $\mathbf{a}_i l_i$ causes the new $(\mathbf{w}^T \mathbf{a}_i) l_i$ to be higher by $\mathbf{a}_i^T \mathbf{a}_i l_i^2 = |\mathbf{a}_i|^2 = 1$. This is good for this $\mathbf{a}_i$. But this change may be bad for other $\mathbf{a}_j$. The proof below shows that this very simple process quickly yields a solution $\mathbf{w}$ provided there exists a solution with a good margin.

Definition 6.1 For a solution $\mathbf{w}$ to (6.2), where $|\mathbf{a}_i| = 1$ for all examples, the margin is defined to be the minimum distance of the hyperplane $\{\mathbf{x} | \mathbf{w}^T \mathbf{x} = 0\}$ to any $\mathbf{a}_i$, namely, $\min_i \frac{(\mathbf{w}^T \mathbf{a}_i) l_i}{|\mathbf{w}|}$. ■

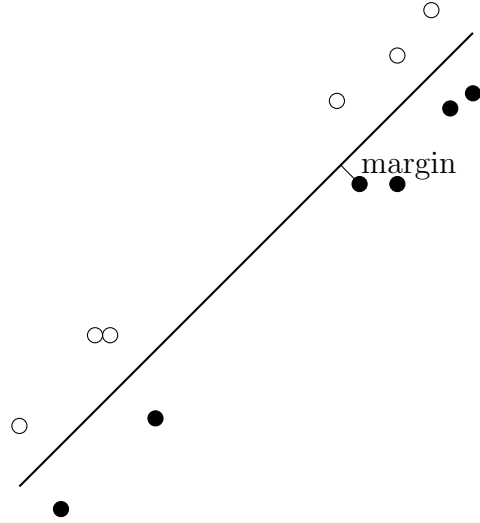


Figure 6.2: Margin of a linear separator.

If we did not require that all $|\mathbf{a}_i| = 1$ in (6.2), then one could artificially increase the margin by scaling up the $\mathbf{a}_i$. If we did not divide by $|\mathbf{w}|$ in the definition of margin, then again, one could artificially increase the margin by scaling $\mathbf{w}$ up. The interesting thing is that the number of steps of the algorithm depends only upon the best margin any solution can achieve, not upon n or d . In practice, the perceptron learning algorithm works well.

Theorem 6.1 *Suppose there is a solution $\mathbf{w}^*$ to (6.2) with margin $\delta > 0$. Then, the perceptron learning algorithm finds some solution $\mathbf{w}$ with $(\mathbf{w}^T \mathbf{a}_i)l_i > 0$ for all i in at most $\frac{1}{\delta^2} - 1$ iterations.*

Proof: Scale $\mathbf{w}^*$ so that $|\mathbf{w}^*| = 1$. Consider the cosine of the angle between the current vector $\mathbf{w}$ and $\mathbf{w}^*$, that is, $\frac{\mathbf{w}^T \mathbf{w}^*}{|\mathbf{w}|}$. In each step of the algorithm, the numerator of this fraction increases by at least δ because

$$(\mathbf{w} + \mathbf{a}_i l_i)^T \mathbf{w}^* = \mathbf{w}^T \mathbf{w}^* + l_i \mathbf{a}_i^T \mathbf{w}^* \geq \mathbf{w}^T \mathbf{w}^* + \delta.$$

On the other hand, the square of the denominator increases by at most one because

$$|\mathbf{w} + \mathbf{a}_i l_i|^2 = (\mathbf{w} + \mathbf{a}_i l_i)^T (\mathbf{w} + \mathbf{a}_i l_i) = |\mathbf{w}|^2 + 2(\mathbf{w}^T \mathbf{a}_i)l_i + |\mathbf{a}_i|^2 l_i^2 \leq |\mathbf{w}|^2 + 1$$

since $\mathbf{w}^T \mathbf{a}_i l_i \leq 0$, the cross term is nonpositive.

After t iterations, $\mathbf{w}^T \mathbf{w}^* \geq (t+1)\delta$ since at the start $\mathbf{w}^T \mathbf{w}^* = l_1(\mathbf{a}_1^T \mathbf{w}^*) \geq \delta$ and at each iteration $\mathbf{w}^T \mathbf{w}^*$ increases by at least δ . Similarly after t iterations $|\mathbf{w}|^2 \leq t+1$ since at the start $|\mathbf{w}| = |\mathbf{a}_1| = 1$ and at each iteration $|\mathbf{w}|^2$ increases by at most one. Thus, the cosine of the angle between $\mathbf{w}$ and $\mathbf{w}^*$ is at least $\frac{(t+1)\delta}{\sqrt{t+1}}$ and the cosine cannot exceed one. Now

$$\frac{(t+1)\delta}{\sqrt{t+1}} \leq 1 \quad \sqrt{t+1}\delta \leq 1 \quad t+1 \leq \frac{1}{\delta^2} \quad t \leq \frac{1}{\delta^2} - 1$$

Therefore the algorithm must stop before $\frac{1}{\delta^2} - 1$ iterations and at termination, $(\mathbf{w}^T \mathbf{a}_i)l_i > 0$ for all i . This proves the theorem. ■

How strong is the assumption that there is a separator with margin at least δ ? Suppose for the moment, the $\mathbf{a}_i$ are picked from the uniform density on the surface of the unit hypersphere. We saw in Chapter 2 that for any fixed hyperplane passing through the origin, most of the mass of the unit sphere is within distance $O(1/\sqrt{d})$ of the hyperplane. So, the probability of one fixed hyperplane having a margin of more than $c/\sqrt{d}$ is low. But this does not mean that there is no hyperplane with a larger margin. By the union bound, one can only assert that the probability of some hyperplane having a large margin is at most the probability of a specific one having a large margin times the number of hyperplanes which is infinite. Later we will see using VC-dimension arguments that indeed the probability of some hyperplane having a large margin is low if the examples are selected at random from the hypersphere. So, the assumption of large margin separators existing may not be valid for the simplest random models. But intuitively, if what is to be learnt, like whether something is a car, is not very hard, then, with enough features in the model, there will not be many “near cars” that could be confused with cars nor many “near non-cars”. In a real problem such as this, uniform density is not a valid assumption. In this case, there should be a large margin separator and the theorem would work.

The question arises as to how small margins can be. Suppose the examples $\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n$ were vectors with d coordinates, each coordinate a 0 or 1 and the decision rule for labeling the examples was the following.

If the first 1 coordinate of the example is odd, label the example +1.

If the first 1 coordinate of the example is even, label the example -1.

This rule can be represented by the decision rule

$$(a_{i1}, a_{i2}, \dots, a_{in}) \left(1, -\frac{1}{2}, \frac{1}{4}, -\frac{1}{8}, \dots\right)^T = a_{i1} - \frac{1}{2}a_{i2} + \frac{1}{4}a_{i3} - \frac{1}{8}a_{i4} + \dots > 0.$$

However, the margin in this example can be exponentially small. Indeed, if for an example $\mathbf{a}$, the first $d/10$ coordinates are all zero, then the margin is $O(2^{-d/10})$.

Maximizing the Margin

In this section, we present an algorithm to find the maximum margin separator. The margin of a solution $\mathbf{w}$ to $(\mathbf{w}^T \mathbf{a}_i)l_i > 0$, $1 \leq i \leq n$, where $|\mathbf{a}_i| = 1$ is $\delta = \min_i \frac{l_i(\mathbf{w}^T \mathbf{a}_i)}{|\mathbf{w}|}$. Since this is not a concave function of $\mathbf{w}$, it is difficult to deal with computationally.

Convex optimization techniques in general can only handle the maximization of concave functions or the minimization of convex functions over convex sets. However, by

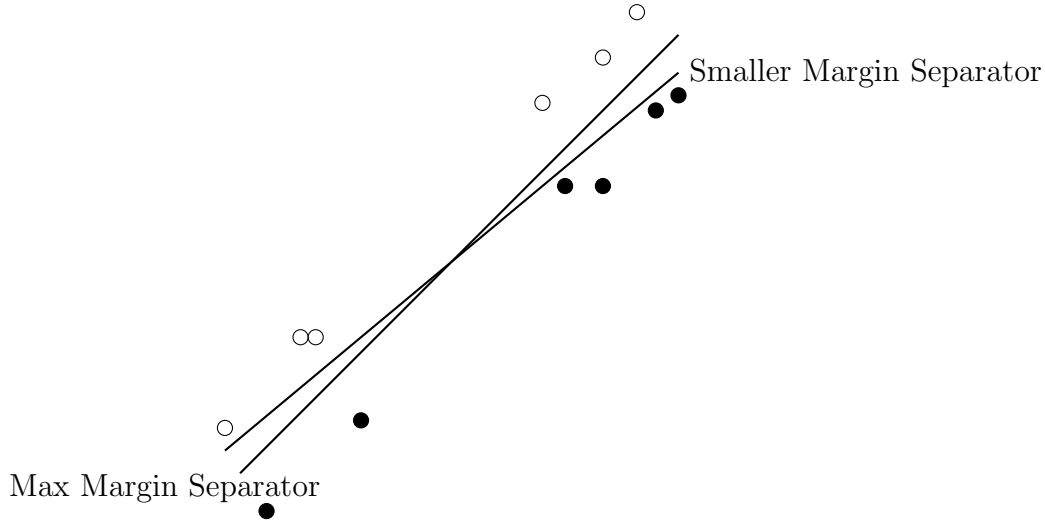


Figure 6.3: Separators with Different Margins

modifying the weight vector, one can convert the optimization problem to one with a concave objective function. Note that

$$l_i \left(\frac{\mathbf{w}^T \mathbf{a}_i}{|\mathbf{w}| \delta} \right) \geq 1$$

for all $\mathbf{a}_i$. Let $\mathbf{v} = \frac{\mathbf{w}}{\delta |\mathbf{w}|}$ be the modified weight vector. Dividing the normalized weight vector $\frac{\mathbf{w}}{|\mathbf{w}|}$ by δ normalizes the margin to one. Maximizing δ is equivalent to minimizing $|\mathbf{v}|$. So the optimization problem is

$$\text{minimize } |\mathbf{v}| \text{ subject to } l_i(\mathbf{v}^T \mathbf{a}_i) > 1, \forall i.$$

Although $|\mathbf{v}|$ is a convex function of the coordinates of $\mathbf{v}$, a better convex function to minimize is $|\mathbf{v}|^2$ since $|\mathbf{v}|^2$ is differentiable. So we reformulate the problem as:

Maximum Margin Problem:

$$\text{minimize } |\mathbf{v}|^2 \text{ subject to } l_i(\mathbf{v}^T \mathbf{a}_i) \geq 1.$$

This convex optimization problem has been much studied and algorithms that use the special structure of this problem solve it more efficiently than general convex optimization methods. We do not discuss these improvements here. An optimal solution $\mathbf{v}$ to this problem has the following property. Let V be the space spanned by the examples $\mathbf{a}_i$ for which there is equality, namely for which $l_i(\mathbf{v}^T \mathbf{a}_i) = 1$. We claim that $\mathbf{v}$ lies in V . If not, $\mathbf{v}$ has a component orthogonal to V . Reducing this component infinitesimally does not violate any inequality, since, we are moving orthogonally to the exactly satisfied constraints; but it does decrease $|\mathbf{v}|$ contradicting the optimality. If V is full dimensional,

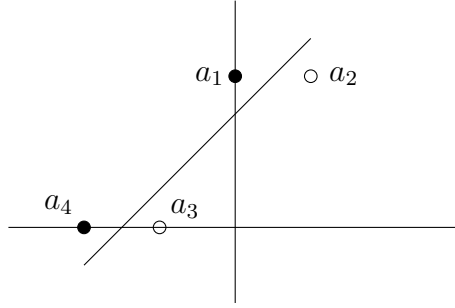


Figure 6.4: The vectors a_1, a_2, a_3 , and a_4 are all support vectors.

then there are d independent examples for which the equality $l_i(\mathbf{v}^T \mathbf{a}_i) = 1$ holds. These d equations then have a unique solution and $\mathbf{v}$ must be that solution. These examples are then called the *support vectors*. The d support vectors determine uniquely the maximum margin separator.

Example: In the example of Figure 6.4 where

$$\begin{array}{cccc} a_1 = [0, 2, 1] & a_2 = [1, 2, 1] & a_3 = [-1, 0, 1] & a_4 = [-2, 0, 1] \\ l_1 = +1 & l_2 = -1 & l_3 = -1 & l_4 = +1 \end{array}$$

all four vectors are support vectors but a_1, a_2 and a_3 are an independent set and determine $[v_1, v_2, v_3]$ uniquely. The solution for v is $v = [-2, 2, -3]$. ■

Linear Separators that classify most examples correctly

It may happen that there are linear separators for which almost all but a small fraction of examples are on the correct side. Going back to (6.2), ask if there is a $\mathbf{w}$ for which at least $(1 - \varepsilon)n$ of the n inequalities in (6.2) are satisfied. Unfortunately, such problems are NP-hard and there are no good algorithms to solve them. A good way to think about this is that we suffer a “loss” of one for each misclassified point and would like to minimize the loss. But this loss function is discontinuous, it goes from 0 to 1 abruptly. However, with a nicer loss function it is possible to solve the problem. One possibility is to introduce slack variables $y_i, i = 1, 2, \dots, n$, where y_i measures how badly the example $\mathbf{a}_i$ is classified. We then include the slack variables in the objective function to be minimized:

$$\begin{array}{ll} \text{minimize} & |\mathbf{v}|^2 + c \sum_{i=1}^n y_i \\ \text{subject to} & \left. \begin{array}{l} (\mathbf{v}^T \mathbf{a}_i) l_i \geq 1 - y_i \\ y_i \geq 0. \end{array} \right\} i = 1, 2, \dots, n \end{array} \quad (6.3)$$

If for some i , $l_i(\mathbf{v}^T \mathbf{a}_i) \geq 1$, then set y_i to its lowest value, namely zero, since each y_i has a positive coefficient in the cost function. If, however, $l_i(\mathbf{v}^T \mathbf{a}_i) < 1$, then set $y_i = 1 - l_i(\mathbf{v}^T \mathbf{a}_i)$,

-1	+1	-1	+1
+1	-1	+1	-1
-1	+1	-1	+1
+1	-1	+1	-1

Figure 6.5: The checker board pattern.

so y_i is just the amount of violation of this inequality. Thus, the objective function is trying to minimize a combination of the total violation as well as $1/\text{margin}$. It is easy to see that this is the same as minimizing

$$|\mathbf{v}|^2 + c \sum_i (1 - l_i(\mathbf{v}^T \mathbf{a}_i))^+,^6$$

subject to the constraints. The second term is the sum of the violations.

6.3 Nonlinear Separators, Support Vector Machines, and Kernels

There are problems where no linear separator exists but where there are nonlinear separators. For example, there may be a polynomial $p(\cdot)$ such that $p(\mathbf{a}_i) > 1$ for all +1 labeled examples and $p(\mathbf{a}_i) < 1$ for all -1 labeled examples. A simple instance of this is the unit square centered at the origin partitioned into four pieces where the top right and the bottom left pieces are the +1 region and the bottom right and the top left are the -1 region. For this, $x_1 x_2 > 0$ for all +1 examples and $x_1 x_2 < 0$ for all -1 examples. So, the polynomial $p(\cdot) = x_1 x_2$ separates the regions. A more complicated instance is the checker-board pattern in Figure 6.5 with alternate +1 and -1 squares.

If we know that there is a polynomial p of degree⁷ at most D such that an example $\mathbf{a}$ has label +1 if and only if $p(\mathbf{a}) > 0$, then the question arises as to how to find such a polynomial. Note that each d -tuple of integers $(i_1, i_2, \dots, i_d)$ with $i_1 + i_2 + \dots + i_d \leq D$ leads to a distinct monomial, $x_1^{i_1} x_2^{i_2} \dots x_d^{i_d}$. So, the number of monomials in the polynomial p is at most the number of ways of inserting $d-1$ dividers into a sequence of $D + d - 1$ positions which is $\binom{D+d-1}{d-1} \leq (D + d - 1)^{d-1}$. Let $m = (D + d - 1)^{d-1}$ be the upper bound

⁶ $x^+ = \begin{cases} 0 & x \leq 0 \\ x & \text{otherwise} \end{cases}$

⁷The degree is the total degree. The degree of a monomial is the sum of the powers of each variable in the monomial and the degree of the polynomial is the maximum degree of its monomials.

on the number of monomials.

By letting the coefficients of the monomials be unknowns, we can formulate a linear program in m variables whose solution gives the required polynomial. Indeed, suppose the polynomial p is

$$p(x_1, x_2, \dots, x_d) = \sum_{\substack{i_1, i_2, \dots, i_d \\ i_1 + i_2 + \dots + i_d \leq D}} w_{i_1, i_2, \dots, i_d} x_1^{i_1} x_2^{i_2} \dots x_d^{i_d}.$$

Then the statement $p(\mathbf{a}_i) > 0$ (recall $\mathbf{a}_i$ is a d -vector) is just a linear inequality in the $w_{i_1, i_2, \dots, i_d}$. However the exponential number of variables for even moderate values of D makes this approach infeasible. Nevertheless, this theoretical approach is useful. First, we clarify the discussion above with an example. Suppose $d = 2$ and $D = 2$. Then the possible (i_1, i_2) form the set $\{(1, 0), (0, 1), (2, 0), (1, 1), (0, 2)\}$. We ought to include the pair $(0, 0)$; but it is convenient to have a separate constant term called b . So we write

$$p(x_1, x_2) = b + w_{10}x_1 + w_{01}x_2 + w_{11}x_1x_2 + w_{20}x_1^2 + w_{02}x_2^2.$$

Each example, $\mathbf{a}_i$, is a 2-vector, (a_{i1}, a_{i2}) . The linear program is

$$\begin{aligned} b + w_{10}a_{i1} + w_{01}a_{i2} + w_{11}a_{i1}a_{i2} + w_{20}a_{i1}^2 + w_{02}a_{i2}^2 &> 0 && \text{if label of } i = +1 \\ b + w_{10}a_{i1} + w_{01}a_{i2} + w_{11}a_{i1}a_{i2} + w_{20}a_{i1}^2 + w_{02}a_{i2}^2 &< 0 && \text{if label of } i = -1. \end{aligned}$$

Note that we “pre-compute” $a_{i1}a_{i2}$, so this does not cause a nonlinearity. The linear inequalities have unknowns that are the $\mathbf{w}$ ’s and b .

The approach above can be thought of as embedding the examples $\mathbf{a}_i$ that are in d -space into a m -dimensional space where there is one coordinate for each $i_1, i_2, \dots, i_d$ summing to at most D , except for $(0, 0, \dots, 0)$, and if $\mathbf{a}_i = (x_1, x_2, \dots, x_d)$, the coordinate is $x_1^{i_1} x_2^{i_2} \dots x_d^{i_d}$. Call this embedding $\boldsymbol{\varphi}(\mathbf{x})$. When $d = D = 2$, as in the above example, $\boldsymbol{\varphi}(\mathbf{x}) = (x_1, x_2, x_1^2, x_1x_2, x_2^2)$. If $d = 3$ and $D = 2$,

$$\boldsymbol{\varphi}(\mathbf{x}) = (x_1, x_2, x_3, x_1x_2, x_1x_3, x_2x_3, x_1^2, x_2^2, x_3^2),$$

and so on. We then try to find a m -dimensional vector $\mathbf{w}$ such that the dot product of $\mathbf{w}$ and $\boldsymbol{\varphi}(\mathbf{a}_i)$ is positive if the label is $+1$ and negative otherwise. Note that this $\mathbf{w}$ is not necessarily the $\boldsymbol{\varphi}$ of some vector in $\mathbf{d}$ space.

Instead of finding any $\mathbf{w}$, we want to find the $\mathbf{w}$ maximizing the margin. As earlier, write this program as

$$\min |\mathbf{w}|^2 \text{ subject to } (\mathbf{w}^T \boldsymbol{\varphi}(\mathbf{a}_i))l_i \geq 1 \text{ for all } i.$$

The major question is whether we can avoid having to explicitly compute the embedding $\boldsymbol{\varphi}$ and the vector $\mathbf{w}$. Indeed, we only need to have $\boldsymbol{\varphi}$ and $\mathbf{w}$ implicitly. This is based on the simple, but crucial observation that any optimal solution $\mathbf{w}$ to the convex program above is a linear combination of the $\boldsymbol{\varphi}(\mathbf{a}_i)$. If $\mathbf{w} = \sum_i y_i \boldsymbol{\varphi}(\mathbf{a}_i)$, then $\mathbf{w}^T \boldsymbol{\varphi}(\mathbf{a}_j)$ can be computed without actually knowing the $\boldsymbol{\varphi}(\mathbf{a}_i)$ but only the products $\boldsymbol{\varphi}(\mathbf{a}_i)^T \boldsymbol{\varphi}(\mathbf{a}_j)$.

Lemma 6.2 *Any optimal solution $\mathbf{w}$ to the convex program above is a linear combination of the $\varphi(\mathbf{a}_i)$.*

Proof: If $\mathbf{w}$ has a component perpendicular to all the $\varphi(\mathbf{a}_i)$, simply zero out that component. This preserves all the inequalities since the $\mathbf{w}^T \varphi(\mathbf{a}_i)$ do not change and decreases $|\mathbf{w}|^2$ contradicting the assumption that $\mathbf{w}$ is an optimal solution. ■

Assume that $\mathbf{w}$ is a linear combination of the $\varphi(\mathbf{a}_i)$. Say $\mathbf{w} = \sum_i y_i \varphi(\mathbf{a}_i)$, where the y_i are real variables. Note that

$$|\mathbf{w}|^2 = \left(\sum_i y_i \varphi(\mathbf{a}_i) \right)^T \left(\sum_j y_j \varphi(\mathbf{a}_j) \right) = \sum_{i,j} y_i y_j \varphi(\mathbf{a}_i)^T \varphi(\mathbf{a}_j).$$

Reformulate the convex program as

$$\text{minimize } \sum_{i,j} y_i y_j \varphi(\mathbf{a}_i)^T \varphi(\mathbf{a}_j) \tag{6.4}$$

$$\text{subject to } l_i \left(\sum_j y_j \varphi(\mathbf{a}_j)^T \varphi(\mathbf{a}_i) \right) \geq 1 \quad \forall i. \tag{6.5}$$

It is important to notice that φ itself is not needed, only the dot products of $\varphi(\mathbf{a}_i)$ and $\varphi(\mathbf{a}_j)$ for all i and j including $i = j$. The *Kernel matrix* K , defined as

$$k_{ij} = \varphi(\mathbf{a}_i)^T \varphi(\mathbf{a}_j),$$

suffices since we can rewrite the convex program as

$$\text{minimize } \sum_{i,j} y_i y_j k_{ij} \quad \text{subject to} \quad l_i \sum_j k_{ij} y_j \geq 1. \tag{6.6}$$

This convex program is called a *support vector machine (SVM)* though it is really not a machine. The advantage is that K has only n^2 entries instead of the $O(d^D)$ entries in each $\varphi(\mathbf{a}_i)$. Instead of specifying $\varphi(\mathbf{a}_i)$, we specify how to get K from the $\mathbf{a}_i$. The specification is usually in closed form. For example, the “Gaussian kernel” is given by:

$$k_{ij} = \varphi(\mathbf{a}_i)^T \varphi(\mathbf{a}_j) = e^{-c|\mathbf{a}_i - \mathbf{a}_j|^2}.$$

We prove shortly that this is indeed a kernel function.

First, an important question arises. Given a matrix K , such as the above matrix for the Gaussian kernel, how do we know that it arises from an embedding φ as the pairwise dot products of the $\varphi(\mathbf{a}_i)$? This is answered in the following lemma.

Lemma 6.3 *A matrix K is a kernel matrix, i.e., there is an embedding φ such that $k_{ij} = \varphi(\mathbf{a}_i)^T \varphi(\mathbf{a}_j)$, if and only if K is positive semidefinite.*

Proof: If K is positive semidefinite, then it can be expressed as $K = BB^T$. Define $\varphi(\mathbf{a}_i)$ to be the i^{th} row of B . Then $k_{ij} = \varphi(\mathbf{a}_i)^T \varphi(\mathbf{a}_j)$. Conversely, if there is an embedding φ such that $k_{ij} = \varphi(\mathbf{a}_i)^T \varphi(\mathbf{a}_j)$, then using the $\varphi(\mathbf{a}_i)$ for the rows of a matrix B , we have that $K = BB^T$ and so K is positive semidefinite. ■

Recall that a function of the form $\sum_{ij} y_i y_j k_{ij} = y^T K y$ is convex if and only if K is positive semidefinite. So the support vector machine problem is a convex program. We may use any positive semidefinite matrix as our kernel matrix.

We now give an important example of a kernel matrix. Consider a set of vectors, $\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n$, and let $k_{ij} = (\mathbf{a}_i^T \mathbf{a}_j)^p$, where p is a positive integer. We prove that the matrix K with elements k_{ij} is positive semidefinite. Suppose $\mathbf{u}$ is any n -vector. We must show that $\mathbf{u}^T K \mathbf{u} = \sum_{ij} k_{ij} u_i u_j \geq 0$.

$$\begin{aligned} \sum_{ij} k_{ij} u_i u_j &= \sum_{ij} u_i u_j (\mathbf{a}_i^T \mathbf{a}_j)^p \\ &= \sum_{ij} u_i u_j \left(\sum_k a_{ik} a_{jk} \right)^p \\ &= \sum_{ij} u_i u_j \left(\sum_{k_1, k_2, \dots, k_p} a_{ik_1} a_{ik_2} \cdots a_{ik_p} a_{jk_1} \cdots a_{jk_p} \right) \text{ by expansion}^8. \end{aligned}$$

Note that $k_1, k_2, \dots, k_p$ need not be distinct. Exchanging the summations and simplifying

$$\begin{aligned} \sum_{ij} u_i u_j \left(\sum_{k_1, k_2, \dots, k_p} a_{ik_1} a_{ik_2} \cdots a_{ik_p} a_{jk_1} \cdots a_{jk_p} \right) &= \sum_{k_1, k_2, \dots, k_p} \sum_{ij} u_i u_j a_{ik_1} a_{ik_2} \cdots a_{ik_p} a_{jk_1} \cdots a_{jk_p} \\ &= \sum_{k_1, k_2, \dots, k_p} \left(\sum_i u_i a_{ik_1} a_{ik_2} \cdots a_{ik_p} \right)^2. \end{aligned}$$

The last term is a sum of squares and thus nonnegative proving that K is positive semidefinite.

From this, it is easy to see that for any set of vectors, $\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n$, and any $c_1, c_2, \dots$ greater than or equal to zero, the matrix K where k_{ij} has an absolutely convergent power

⁸Here a_{ik} denotes the k^{th} coordinate of the vector $\mathbf{a}_i$

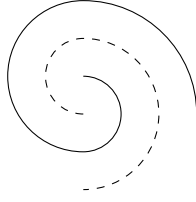


Figure 6.6: Two curves

series expansion $k_{ij} = \sum_{p=0}^{\infty} c_p (\mathbf{a}_i^T \mathbf{a}_j)^p$ is positive semidefinite. For any $\mathbf{u}$,

$$\mathbf{u}^T K \mathbf{u} = \sum_{ij} u_i k_{ij} u_j = \sum_{ij} u_i \left(\sum_{p=0}^{\infty} c_p (\mathbf{a}_i^T \mathbf{a}_j)^p \right) u_j = \sum_{p=0}^{\infty} c_p \sum_{i,j} u_i u_j (\mathbf{a}_i^T \mathbf{a}_j)^p \geq 0.$$

Lemma 6.4 For any set of vectors, $\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n$, the matrix K given by $k_{ij} = e^{-|\mathbf{a}_i - \mathbf{a}_j|^2 / (2\sigma^2)}$ is positive semidefinite for any value of σ .

Proof:

$$e^{-|\mathbf{a}_i - \mathbf{a}_j|^2 / 2\sigma^2} = e^{-|\mathbf{a}_i|^2 / 2\sigma^2} e^{-|\mathbf{a}_j|^2 / 2\sigma^2} e^{\mathbf{a}_i^T \mathbf{a}_j / \sigma^2} = \left(e^{-|\mathbf{a}_i|^2 / 2\sigma^2} e^{-|\mathbf{a}_j|^2 / 2\sigma^2} \right) \sum_{t=0}^{\infty} \left(\frac{(\mathbf{a}_i^T \mathbf{a}_j)^t}{t! \sigma^{2t}} \right).$$

The matrix L given by $l_{ij} = \sum_{t=0}^{\infty} \left(\frac{(\mathbf{a}_i^T \mathbf{a}_j)^t}{t! \sigma^{2t}} \right)$ FF is positive semi-definite. Now K can be written as DLD^T , where D is the diagonal matrix with $e^{-|\mathbf{a}_i|^2 / 2\sigma^2}$ as its $(i, i)^{th}$ entry. So K is positive semi-definite. ■

Example: (Use of the Gaussian Kernel) Consider a situation where examples are points in the plane on two juxtaposed curves, the solid curve and the dotted curve shown in Figure 6.6, where points on the first curve are labeled +1 and points on the second curve are labeled -1. Suppose examples are spaced δ apart on each curve and the minimum distance between the two curves is $\Delta \gg \delta$. Clearly, there is no half-space in the plane that classifies the examples correctly. Since the curves intertwine a lot, intuitively, any polynomial which classifies them correctly must be of high degree. Consider the Gaussian kernel $e^{-|\mathbf{a}_i - \mathbf{a}_j|^2 / \delta^2}$. For this kernel, the K has $k_{ij} \approx 1$ for adjacent points on the same curve and $k_{ij} \approx 0$ for all other pairs of points. Reorder the examples, first listing in order all examples on the solid curve, then on the dotted curve. K has the block form: $K = \begin{pmatrix} K_1 & 0 \\ 0 & K_2 \end{pmatrix}$, where K_1 and K_2 are both roughly the same size and are both block matrices with 1's on the diagonal and slightly smaller constants on the diagonals one off from the main diagonal and then exponentially falling off with distance from the diagonal.

The SVM is easily seen to be essentially of the form:

$$\begin{aligned} & \text{minimize} && \mathbf{y}_1^T K_1 \mathbf{y}_1 + \mathbf{y}_2^T K_2 \mathbf{y}_2 \\ & \text{subject to} && K_1 \mathbf{y}_1 \geq 1 \quad \text{and} \quad K_2 \mathbf{y}_2 \leq -1. \end{aligned}$$

This separates into two programs, one for $\mathbf{y}_1$ and the other for $\mathbf{y}_2$. From the fact that $K_1 = K_2$, the solution will have $\mathbf{y}_2 = -\mathbf{y}_1$. Further by the structure which is essentially the same everywhere except at the ends of the curves, the entries in $\mathbf{y}_1$ will all be essentially the same as will the entries in $\mathbf{y}_2$. Thus, the entries in $\mathbf{y}_1$ will be 1 everywhere and the entries in $\mathbf{y}_2$ will be -1 everywhere. Let l_i be the ± 1 labels for the points. The y_i values provide a nice simple classifier, namely $l_i y_i > 1$. ■

6.4 Strong and Weak Learning - Boosting

A strong learner is an algorithm that takes n labeled examples and produces a classifier that correctly labels each of the given examples. Since the learner is given the n examples with their labels and is responsible only for the given training examples, it seems a trivial task. Just store the examples and labels in a table and each time we are asked for the label of one of the examples, do a table look-up. By Occam's razor principle, the classifier produced by the learner should be considerably more concise than a table of the given examples. The time taken by the learner, and the length/complexity of the classifier output are both parameters by which we measure the learner. For now we focus on a different aspect. The word strong refers to the fact that the output classifier must label all the given examples correctly; no errors are allowed.

A weak learner is allowed to make mistakes. It is only required to get a strict majority, namely, a $(\frac{1}{2} + \gamma)$ fraction of the examples correct where γ is a positive real number. This seems very weak. But with a slight generalization of weak learning and using a technique called boosting, strong learning can be accomplished with a weak learner.

Definition 6.2 (*Weak learner*) Suppose $U = \{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n\}$ are n labeled examples. A weak learner is an algorithm that given the examples, their labels, and a nonnegative real weight w_i on each example $\mathbf{a}_i$ as input, produces a classifier that correctly labels a subset of examples with total weight at least $(\frac{1}{2} + \gamma) \sum_{i=1}^n w_i$. ■

A strong learner can be built by making $O(\log n)$ calls to a weak learner by a method called boosting. Boosting makes use of the intuitive notion that if an example was misclassified, one needs to pay more attention to it.

Example: Illustration of boosting

$\begin{array}{ccc} x & x & x \\ 0 & 0 & x \\ 0 & 0 & x \end{array}$	$\begin{array}{cc c} x & x & x \\ 0 & 0 & x \\ 0 & 0 & x \end{array}$	$\begin{array}{ccc} x & x & x \\ 0 & 0 & x \\ 0 & 0 & x \end{array}$	$\begin{array}{ccc} x & x & x \\ 0 & 0 & x \\ 0 & 0 & x \end{array}$	$\begin{array}{ccc} x & x & x \\ 0 & 0 & x \\ 0 & 0 & x \end{array}$
--	---	--	--	--

Learn x 's from 0's. Items above or to the right of a line are classified as x 's.

$\begin{array}{ccc} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{array}$	$\begin{array}{ccc} 1 & 1 & 1 \\ 1 & 1 & 1+\epsilon \\ 1 & 1 & 1+\epsilon \end{array}$	$\begin{array}{ccc} 1+\epsilon & 1+\epsilon & 1 \\ 1 & 1 & 1+\epsilon \\ 1 & 1 & 1+\epsilon \end{array}$	$\begin{array}{ccc} 1+\epsilon & 1+\epsilon & 1 \\ 1 & 1+\epsilon & 1+\epsilon \\ 1 & 1 & 1+\epsilon \end{array}$	$\begin{array}{ccc} (1+\epsilon)2 & 1+\epsilon & 1 \\ 1 & 1+\epsilon & 1+\epsilon \\ 1 & 1 & 1+\epsilon \end{array}$
--	--	--	---	--

Weight of each example over time

$\begin{array}{ccc} 0 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 1 \end{array}$	$\begin{array}{ccc} 1 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 1 \end{array}$	$\begin{array}{ccc} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{array}$	$\begin{array}{ccc} 2 & 1 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{array}$	$\begin{array}{ccc} 2 & 1 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{array}$
--	--	--	--	--

Number of times misclassified

The top row indicates the results of the weak learner and the middle row indicates the weight applied to each example. In the first application, the weak learner misclassified the bottom two elements of the rightmost column. Thus, in the next application the weights of these two items were increased from 1 to $1 + \epsilon$. The bottom row of matrices indicates how often each element was misclassified. Since no element was misclassified more than two out of five times, the results of labeling each example by the way it was classified a majority of times give the correct labeling of all elements. ■

Boosting algorithm

Make the first call to the weak learner with all w_i set equal to 1.

At time $t + 1$ multiply the weight of each example that was misclassified the previous time by $1 + \epsilon$. Leave the other weights as they are. Make a call to the weak learner.

After T steps, stop and output the following classifier:

Label each of the examples $\{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n\}$ by the label given to it by a majority of calls to the weak learner. Assume T is odd, so there is no tie for the majority.

■

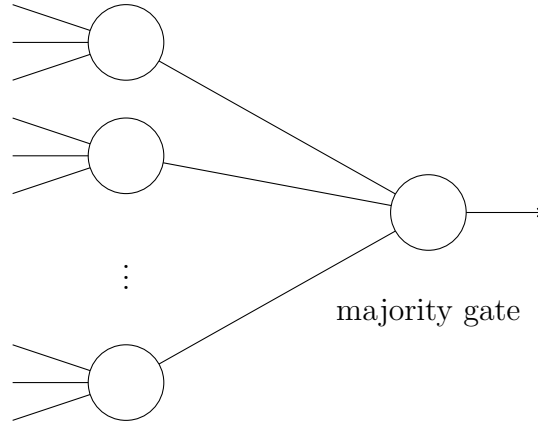


Figure 6.7: Learner produced by boosting algorithm

Suppose m is the number of examples the final classifier gets wrong. Each of these m examples was misclassified at least $T/2$ times so each has weight at least $(1 + \varepsilon)^{T/2}$. This says the total weight is at least $m(1 + \varepsilon)^{T/2}$. On the other hand, at time $t + 1$, only the weight of examples misclassified at time t were increased. By the property of weak learning, the total weight of misclassified examples is at most $f = (\frac{1}{2} - \gamma)$ of the total weight at time t . Let $\text{weight}(t)$ be the total weight at time t . Then

$$\begin{aligned} \text{weight}(t + 1) &\leq ((1 + \varepsilon)f + 1 - f) \times \text{weight}(t) \\ &= (1 + \varepsilon f) \times \text{weight}(t) \\ &\leq \left(1 + \frac{\varepsilon}{2} - \gamma\varepsilon\right) \times \text{weight}(t). \end{aligned}$$

Thus, since $\text{weight}(0) = n$,

$$m(1 + \varepsilon)^{T/2} \leq \text{Total weight at end} \leq n(1 + \frac{\varepsilon}{2} - \gamma\varepsilon)^T.$$

Taking logarithms

$$\ln m + \frac{T}{2} \ln(1 + \varepsilon) \leq \ln n + T \ln(1 + \frac{\varepsilon}{2} - \gamma\varepsilon).$$

To a first order approximation, $\ln(1 + \delta) \approx \delta$ for δ small. Make ε , the amount the weights of misclassified examples are increased by, a small constant, say $\varepsilon = 0.01$. Then $\ln m \leq \ln n - T\gamma\varepsilon$. Let $T = (1 + \ln n)/\gamma\varepsilon$. Then $\ln m \leq -1$ and $m \leq \frac{1}{e}$. Thus, the number of misclassified items, m , is less than one and hence must be zero.

6.5 Number of Examples Needed for Prediction: VC-Dimension Training and Prediction

Up to this point, we dealt only with training examples and focused on building a classifier that works correctly on them. Of course, the ultimate purpose is prediction of labels on future examples. In the car versus non-car example, we want our classifier to classify future feature vectors as car or non-car without human input. Clearly, we cannot expect the classifier to predict every example correctly. To measure how good the classifier is, we attach a probability distribution on the space of examples and measure the probability of misclassification. The reason for attaching a probability distribution is that we want the classifier to correctly classify likely examples but are not so concerned about examples that almost never arise.

A second question is how many training examples suffice so that as long as a classifier gets all the training examples correct (strong learner), the probability that it makes a prediction error (measured with the same probability distribution as used to select training examples) of more than ε is less than δ ? Ideally, we would like this number to be sufficient whatever the unknown probability distribution is. The theory of VC-dimension will provide an answer to this.

A Sampling Motivation

The concept of VC-dimension is fundamental and is the backbone of learning theory. It is also useful in many other contexts. Our first motivation will be from a database example. Consider a database consisting of the salary and age of each employee in a company and a set of queries of the form: how many individuals between ages 35 and 45 have a salary between \$60,000 and \$70,000? Each employee is represented by a point in the plane where the coordinates are age and salary. The query asks how many data points fall within an axis-parallel rectangle. One might want to select a fixed sample of the data before queries arrive and estimate the number of points in a query rectangle based on the number of sample points in the rectangle. For one rectangle the probability that the estimate is off by more than an ϵ -fraction can be made less than δ by making the sample large enough. However, we want the sample to work for all rectangles. At first, such an estimate would not seem to work. Applying a union bound, the probability that there exists a rectangle that the sample fails to work for is at most the product of the probability that the sample fails for one particular rectangle times the number of possible rectangles. But, there are an infinite number of possible rectangles. So such a simple union bound argument does not give a finite upper bound on the probability that the sample fails to work for rectangles.

Define two axis-parallel rectangles to be equivalent if they contain the same data points. If there are n data points, only $O(n^4)$ of the 2^n subsets can correspond to the set of points in a rectangle. To see this, consider any rectangle R . If one of its sides does not pass through one of the n points that is inside the rectangle, then move the side parallel to itself until for the first time it passes through one of the n points inside the rectangle. Clearly, the set of points in R and the new rectangle are the same since the edge did not

“cross” any point. By a similar process, modify all four sides, so that there is at least one point on each side of the rectangle. Now, the number of rectangles with at least one point on each side is at most $O(n^4)$. The exponent four plays an important role; it will turn out to be the VC-dimension of axis-parallel rectangles.

Let U be a set of n points in the plane where each point corresponds to one employee’s age and salary. Let $\varepsilon > 0$ be a given error parameter. Pick a random sample S of size s from U . Given a query rectangle R , estimate $|R \cap U|$ by the quantity $\frac{n}{s}|R \cap S|$. This is the number of employees in the sample within the ranges scaled up by $\frac{n}{s}$, since we picked a sample of size s out of n . We wish to assert that the fractional error for a random sample of size s is at most ε for every rectangle R , i.e., that

$$\left| |R \cap U| - \frac{n}{s}|R \cap S| \right| \leq \varepsilon n$$

for every R . Of course, the assertion is not absolute, there is a small probability that the sample is atypical, for example picking no points from a rectangle R which has a lot of points. We can only assert the above with high probability or that its negation holds with very low probability. That is,

$$\text{Prob} \left(\exists \text{ an } R \left| |R \cap U| - \frac{n}{s}|R \cap S| \right| > \varepsilon n \right) \leq \delta, \quad (6.7)$$

where $\delta > 0$ is another error parameter. Note that it is very important that our sample S be good for every possible query, since we do not know beforehand which queries will arise.

How many samples are necessary to ensure that (6.7) holds? Pick s samples uniformly at random from the n points in U . For one fixed R , the number of samples in R is a random variable that is the sum of s independent 0-1 random variables, each with probability $q = \frac{|R \cap U|}{n}$ of having value one. The distribution of $|R \cap S|$ is Binomial(s, q). Using Chernoff bounds, for $0 \leq \varepsilon \leq 1$,

$$\text{Prob} \left(\left| |R \cap U| - \frac{n}{s}|R \cap S| \right| > \varepsilon n \right) \leq 2e^{-\varepsilon^2 s / (3q)} \leq 2e^{-\varepsilon^2 s / 3}.$$

Using the union bound and noting that there are only $O(n^4)$ possible sets $R \cap U$ yields

$$\text{Prob} \left(\exists \text{ an } R \left| |R \cap U| - \frac{n}{s}|R \cap S| \right| > \varepsilon n \right) \leq cn^4 e^{-\varepsilon^2 s / 3}$$

for some sufficiently large c . Setting

$$s \geq \frac{3}{\varepsilon^2} \left(5 \ln n + \ln \frac{1}{\delta} \right)$$

ensures (6.7) when n is sufficiently large. In fact, we will see later that even the logarithmic dependence on n can be avoided. As long as s is at least a certain number depending only upon the error ε and the VC-dimension of the set of shapes, (6.7) will hold.

In another situation, suppose we have an unknown probability distribution p over the plane and ask what is the probability mass $p(R)$ of a query rectangle R ? We might estimate the probability mass by first drawing a sample S of size s in s independent trials, each draw according to p , and wish to know how far the sample estimate $|S \cap R|/s$ is from the probability mass $p(R)$. Again, we would like the estimate to be good for every rectangle. This is a more general problem than the first problem of estimating $|R \cap U|$. The first problem is the particular case where U consists of n points in the plane and the probability distribution p has value $\frac{1}{n}$ at each of n points. Then $\frac{1}{n}|R \cap U| = p(R)$.

There is no simple argument bounding the number of rectangles to $O(n^4)$ for the general problem. Moving the sides of the rectangle is no longer valid, since it could change the enclosed probability mass. Further, p could be a continuous distribution, where the analog of n would be infinite. So the argument above using the union bound would not solve the problem. The VC-dimension argument will yield the desired result for the more general situation.

The question is also of interest for shapes other than rectangles. Indeed, half-spaces in d -dimensions is an important class of “shapes”, since they correspond to threshold gates. A class of regions such as halfspaces or rectangles has a parameter called VC-dimension and we can bound the probability of the discrepancy between the sample estimate and the probability mass in terms of the VC-dimension of the shapes allowed. That is,

$$|\text{prob mass} - \text{estimate}| < \varepsilon$$

with probability $1 - \delta$ where δ depends on ε and the VC-dimension.

In summary, we would like to create a sample of the data base without knowing which query we will face, knowing only the family of possible queries such as rectangles. We would like our sample to work well for every possible query from the class. With this motivation, we introduce VC-dimension and later relate it to learning.

6.6 Vapnik-Chervonenkis or VC-Dimension

A *set system* $(U, \mathbb{S})$ consists of a set U along with a collection $\mathbb{S}$ of subsets of U . The set U may be finite or infinite. An example of a set system is the set $U = R^2$ of points in the plane, with $\mathbb{S}$ being the collection of all axis-parallel rectangles. Each rectangle is viewed as the set of points in it.

Let $(U, \mathbb{S})$ be a set system. A subset $A \subseteq U$ is *shattered* by $\mathbb{S}$ if each subset of A can be expressed as the intersection of an element of $\mathbb{S}$ with A . The VC-dimension of the set system $(U, \mathbb{S})$ is the maximum size of any subset of U shattered by $\mathbb{S}$.

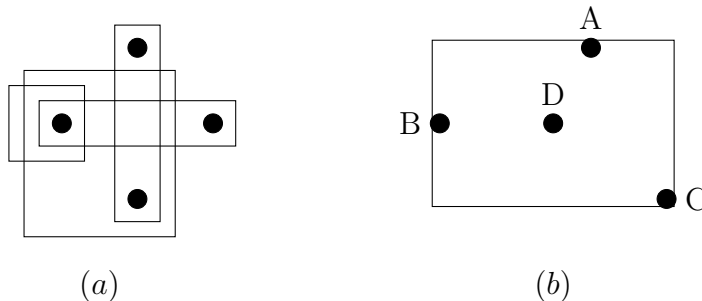


Figure 6.8: (a) shows a set of four points that can be shattered by rectangles along with some of the rectangles that shatter the set. Not every set of four points can be shattered as seen in (b). Any rectangle containing points A, B, and C must contain D. No set of five points can be shattered by rectangles with axis-parallel edges. No set of three collinear points can be shattered, since any rectangle that contains the two end points must also contain the middle point. More generally, since rectangles are convex, a set with one point inside the convex hull of the others cannot be shattered.

6.6.1 Examples of Set Systems and Their VC-Dimension

Rectangles with axis-parallel edges

There exist sets of four points that can be shattered by rectangles with axis-parallel edges. For example, four points at the vertices of a diamond. However, rectangles with axis-parallel edges cannot shatter any set of five points. To see this, assume for contradiction that there is a set of five points shattered by the family of axis-parallel rectangles. Find the minimum enclosing rectangle for the five points. For each edge there is at least one point that has stopped its movement. Identify one such point for each edge. The same point maybe identified as stopping two edges if it is at a corner of the minimum enclosing rectangle. If two or more points have stopped an edge, designate only one as having stopped the edge. Now, at most four points have been designated. Any rectangle enclosing the designated points must include the undesigned points. Thus, the subset of designated points cannot be expressed as the intersection of a rectangle with the five points. Therefore, the VC-dimension of axis-parallel rectangles is four.

Intervals of the reals

Intervals on the real line can shatter any set of two points but no set of three points since the subset of the first and last points cannot be isolated. Thus, the VC-dimension of intervals is two.

Pairs of intervals of the reals

Consider the family of pairs of intervals, where a pair of intervals is viewed as the set of points that are in at least one of the intervals, in other words, their set union. There exists a set of size four that can be shattered but no set of size five since the subset of first, third, and last point cannot be isolated. Thus, the VC-dimension of pairs of intervals is four.

Convex polygons

Consider the set system of all convex polygons in the plane. For any positive integer n , place n points on the unit circle. Any subset of the points are the vertices of a convex polygon. Clearly that polygon will not contain any of the points not in the subset. This shows that convex polygons can shatter arbitrarily large sets, so the VC-dimension is infinite.

Half spaces in d -dimensions

Define a half space to be the set of all points on one side of a hyper plane, i.e., a set of the form $\{\mathbf{x} | \mathbf{a}^T \mathbf{x} \geq a_0\}$. The VC-dimension of half spaces in d -dimensions is $d + 1$.

There exists a set of size $d + 1$ that can be shattered by half spaces. Select the d unit-coordinate vectors plus the origin to be the $d + 1$ points. Suppose A is any subset of these $d + 1$ points. Without loss of generality assume that the origin is in A . Take a 0-1 vector $\mathbf{a}$ which has 1's precisely in the coordinates corresponding to vectors not in A . Clearly A lies in the half-space $\mathbf{a}^T \mathbf{x} \leq 0$ and the complement of A lies in the complementary half-space.

We now show that no set of $d + 2$ points can be shattered by half spaces. To this end, we first show that any set of $d + 2$ points can be partitioned into two disjoint subsets A and B of points whose convex hulls intersect. Let $\text{convex}(A)$ and $\text{convex}(B)$ denote the convex hull of the sets of points in A and B . First consider four points in 2-dimensions. If any three of the points lie on a straight line, then the mid point lies in the convex hull of the other two. Thus, assume that no three of the points lie on a straight line. Select three of the points. The three points must form a triangle. Extend the edges of the triangle to infinity. The three lines divide the plane into seven regions, one finite and six infinite. Place the fourth point in the plane. If the point is placed in the triangle, then it and the convex hull of the triangle intersect. If the fourth point lies in a two sided infinite region, the convex hull of the point plus the two opposite points of the triangle contains the third vertex of the triangle. If the fourth point is in a three sided region, the convex hull of the point plus the opposite vertex of the triangle intersects the convex hull of the other two points of the triangle.

Consider $d + 2$ points in d -dimensions and assume we have established the claim for dimensions less than d . Thus, if $d + 1$ points lie on a $d - 1$ -dimensional hyper plane, then we

are done. Assume that $d+1$ points are in general position and form a hyper tetrahedron in d -space. Extend the $d-1$ dimensional faces to hyper planes. The hyper planes partition d -space into a finite region, the tetrahedron, and a number of infinite regions. Each infinite region contain a vertex, edge, face, etc. of the finite region. Refer to the component of the finite region that meets an infinite region as a face of the tetrahedron. Let the points of the face be one subset and the remaining vertices of the tetrahedron plus a point in the infinite region be the other subset. The convex hulls of these two subsets intersect. The reader is encouraged to develop these ideas into a geometric proof. Here, instead, we present an algebraic proof.

Theorem 6.5 (Radon): *Any set $S \subseteq R^d$ with $|S| \geq d+2$, can be partitioned into two disjoint subsets A and B such that $\text{convex}(A) \cap \text{convex}(B) \neq \emptyset$.*

Proof: Without loss of generality, assume $|S| = d+2$. Form a $d \times (d+2)$ matrix with one column for each point of S . Call the matrix A . Add an extra row of all 1's to construct a $(d+1) \times (d+2)$ matrix B . Clearly, since the rank of this matrix is at most $d+1$, the columns are linearly dependent. Say $\mathbf{x} = (x_1, x_2, \dots, x_{d+2})$ is a nonzero vector with $B\mathbf{x} = 0$. Reorder the columns so that $x_1, x_2, \dots, x_s \geq 0$ and $x_{s+1}, x_{s+2}, \dots, x_{d+2} < 0$. Normalize $\mathbf{x}$ so $\sum_{i=1}^s |x_i| = 1$. Let $\mathbf{b}_i$ (respectively $\mathbf{a}_i$) be the i^{th} column of B (respectively A). Then, $\sum_{i=1}^s |x_i| \mathbf{b}_i = \sum_{i=s+1}^{d+2} |x_i| \mathbf{b}_i$ from which it follows that $\sum_{i=1}^s |x_i| \mathbf{a}_i = \sum_{i=s+1}^{d+2} |x_i| \mathbf{a}_i$ and $\sum_{i=1}^s |x_i| = \sum_{i=s+1}^{d+2} |x_i|$. Since $\sum_{i=1}^s |x_i| = 1$ and $\sum_{i=s+1}^{d+2} |x_i| = 1$ each side of $\sum_{i=1}^s |x_i| \mathbf{a}_i = \sum_{i=s+1}^{d+2} |x_i| \mathbf{a}_i$ is a convex combination of columns of A which proves the theorem. Thus, S can be partitioned into two sets, the first consisting of the first s points after the rearrangement and the second consisting of points $s+1$ through $d+2$. Their convex hulls intersect as required. ■

Radon's theorem immediately implies that half-spaces in d -dimensions do not shatter any set of $d+2$ points. Divide the set of $d+2$ points into sets A and B where $\text{convex}(A) \cap \text{convex}(B) \neq \emptyset$. Suppose that some half space separates A from B . Then the half space separates the convex hulls of A and B . Thus, $\text{convex}(A) \cap \text{convex}(B) = \emptyset$ a contradiction. Therefore, no set of $d+2$ points can be shattered by half planes in d -dimensions.

Spheres in d -dimensions

A *sphere* in d -dimensions is a set of points of the form $\{\mathbf{x} \mid |\mathbf{x} - \mathbf{x}_0| \leq r\}$. The VC-dimension of spheres is $d+1$. It is the same as that of half spaces. First, we prove that no set of $d+2$ points can be shattered by spheres. Suppose some set S with $d+2$ points can be shattered. Then for any partition A_1 and A_2 of S , there are spheres B_1 and B_2 such that $B_1 \cap S = A_1$ and $B_2 \cap S = A_2$. Now B_1 and B_2 may intersect, but there is no

point of S in their intersection. It is easy to see that there is a hyperplane perpendicular to the line joining the centers of the two spheres with all of A_1 on one side and all of A_2 on the other and this implies that half spaces shatter S , a contradiction. Therefore no $d + 2$ points can be shattered by hyperspheres.

It is also not difficult to see that the set of $d + 1$ points consisting of the unit-coordinate vectors and the origin can be shattered by spheres. Suppose A is a subset of the $d + 1$ points. Let a be the number of unit vectors in A . The center $\mathbf{a}_0$ of our sphere will be the sum of the vectors in A . For every unit vector in A , its distance to this center will be $\sqrt{a - 1}$ and for every unit vector outside A , its distance to this center will be $\sqrt{a + 1}$. The distance of the origin to the center is $\sqrt{a}$. Thus, we can choose the radius so that precisely the points in A are in the hypersphere.

Finite sets

The system of finite sets of real numbers can shatter any finite set of real numbers and thus the VC-dimension of finite sets is infinite.

6.6.2 The Shatter Function

Consider a set system $(U, \mathbb{S})$ of finite VC-dimension d . For $n \leq d$ there exists a subset $A \subseteq U$, $|A| = n$, such that A can be shattered into 2^n pieces. This raises the question for $|A| = n$, $n > d$, as to what is the maximum number of subsets of A expressible as $S \cap A$ for $S \in \mathbb{S}$. We shall see that this maximum number is at most a polynomial in n with degree d .

The *shatter function* $\pi_{\mathbb{S}}(n)$ of a set system $(U, \mathbb{S})$ is the maximum number of subsets that can be defined by the intersection of sets in $\mathbb{S}$ with some n element subset A of U . Thus

$$\pi_{\mathbb{S}}(n) = \max_{\substack{A \subseteq U \\ |A|=n}} |\{A \cap S \mid S \in \mathbb{S}\}|$$

For small values of n , $\pi_{\mathbb{S}}(n)$ will grow as 2^n . Once n equals the VC-dimension of $\mathbb{S}$, it grows more slowly. The definition of VC-dimension can clearly be reformulated as $\dim(\mathbb{S}) = \max\{n \mid \pi_{\mathbb{S}}(n) = 2^n\}$. Curiously, the growth of $\pi_{\mathbb{S}}(n)$ must be either polynomial or exponential in n . If the growth is exponential, then the VC-dimension of $\mathbb{S}$ is infinite.

Examples of set systems and their shatter function.

Example: Half spaces and circles in the plane have VC-dimension three. So, their shatter function is 2^n for $n=1, 2$, and 3 . For $n > 3$, their shatter function grows as a polynomial of degree three in n . Axis-parallel rectangles have VC-dimension four and thus their shatter function is 2^n for $n=1, 2, 3$, and 4 . For $n > 4$, their shatter function grows as a polynomial of degree four in n .

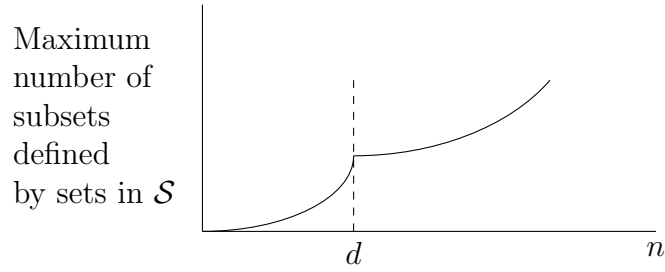


Figure 6.9: The shatter function for a set system of VC-dimension d

We already saw that for axis-parallel rectangles in the plane, there are at most $O(n^4)$ possible subsets of an n element set that arise as intersections with rectangles. The argument was that one can move the sides of the rectangle until each side is “blocked” by one point. We also saw that the VC-dimension of axis-parallel rectangles is four. We will see here that the two fours, one in the exponent of n and the other the VC-dimension, being equal is no accident. There is another four related to rectangles, that is, it takes four parameters to specify an axis-parallel rectangle. Although the VC-dimension of a collection of sets is often closely related to the number of free parameters, this latter four is a coincidence.

6.6.3 Shatter Function for Set Systems of Bounded VC-Dimension

For any set system $(U, \mathbb{S})$ of VC-dimension d , the quantity

$$\sum_{i=0}^d \binom{n}{i} = \binom{n}{0} + \binom{n}{1} + \cdots + \binom{n}{d} \leq 2n^d$$

bounds the shatter function $\pi_{\mathbb{S}}(n)$. That is, $\sum_{i=0}^d \binom{n}{i}$ bounds the number of subsets of any n point subset of U that can be expressed as the intersection with a set of $\mathbb{S}$. Thus, the shatter function $\pi_{\mathbb{S}}(n)$ is either 2^n if d is infinite or it is bounded by a polynomial of degree d .

Lemma 6.6 *For any set system $(U, \mathbb{S})$ of VC-dimension at most d , $\pi_{\mathbb{S}}(n) \leq \sum_{i=0}^d \binom{n}{i}$ for all n .*

Proof: The proof is by induction on d and n . The base case will handle all pairs (d, n) with either $n \leq d$ or $d = 0$. The general case (d, n) will use the inductive assumption on the cases $(d - 1, n - 1)$ and $(d, n - 1)$.

For $n \leq d$, $\sum_{i=0}^d \binom{n}{i} = \sum_{i=0}^n \binom{n}{i} = 2^n$ and $\pi_{\mathbb{S}}(n) = 2^n$. For $d = 0$, a set system $(U, \mathbb{S})$ can have at most one set in $\mathbb{S}$ since if there were two sets in $\mathbb{S}$ there would exist a set A consisting of a single element that was contained in one of the sets but not in the other that could

be shattered. If $\mathbb{S}$ contains only one set, then $\pi_{\mathbb{S}}(n) = 1$ for all n and for $d = 0$, $\sum_{i=0}^d \binom{n}{i} = 1$.

Consider the case for general d and n . Select a subset A of U of size n such that $\pi_{\mathbb{S}}(n)$ subsets of A can be expressed as $A \cap S$ for S in $\mathbb{S}$. Without loss of generality, we may assume that $U = A$ and replace each set $S \in \mathbb{S}$ by $S \cap A$ removing duplicate sets; i.e., if $S_1 \cap A = S_2 \cap A$ for S_1 and S_2 in $\mathbb{S}$, keep only one of them. Now each set in $\mathbb{S}$ corresponds to a subset of A and $\pi_{\mathbb{S}}(n) = |\mathbb{S}|$. Thus, to show $\pi_{\mathbb{S}}(n) \leq \sum_{i=0}^d \binom{n}{i}$, we only need to show

$$|\mathbb{S}| \leq \sum_{i=0}^d \binom{n}{i}.$$

Remove some element u from the set A and from each set in $\mathbb{S}$. Consider the set system $\mathbb{S}_1 = (A - \{u\}, \{S - \{u\} | S \in \mathbb{S}\})$. For $S \subseteq A - \{u\}$, if exactly one of S and $S \cup \{u\}$ is in $\mathbb{S}$, then the set S contributes one set to both $\mathbb{S}$ and $\mathbb{S}_1$, whereas, if both S and $S \cup \{u\}$ are in $\mathbb{S}$, then they together contribute two sets to $\mathbb{S}$, but only one to $\mathbb{S}_1$. Thus $|\mathbb{S}_1|$ is less than $|\mathbb{S}|$ by the number of pairs of sets in $\mathbb{S}$ that differ only in the element u . To account for this difference, define another set system

$$\mathbb{S}_2 = (A - \{u\}, \{S | \text{both } S \text{ and } S \cup \{u\} \text{ are in } \mathbb{S}\}).$$

Then

$$|\mathbb{S}| = |\mathbb{S}_1| + |\mathbb{S}_2| = \pi_{\mathbb{S}_1}(n-1) + \pi_{\mathbb{S}_2}(n-1)$$

or

$$\pi_{\mathbb{S}}(n) = \pi_{\mathbb{S}_1}(n-1) + \pi_{\mathbb{S}_2}(n-1).$$

We make use of two facts

(1) $\mathbb{S}_1$ has dimension at most d , and

(2) $\mathbb{S}_2$ has dimension at most $d-1$.

(1) follows because if $\mathbb{S}_1$ shatters a set of cardinality $d+1$, then $\mathbb{S}$ also would shatter that set producing a contradiction. (2) follows because if $\mathbb{S}_2$ shattered a set $B \subseteq A - \{u\}$ with $|B| \geq d$, then $B \cup \{u\}$ would be shattered by S where $|B \cup \{u\}| \geq d+1$, again producing a contradiction.

By the induction hypothesis applied to $\mathbb{S}_1$, we have $|\mathbb{S}_1| = \pi_{\mathbb{S}_1}(n-1) \leq \sum_{i=0}^d \binom{n-1}{i}$. By the induction hypotheses applied to $\mathbb{S}_2$, we have $|\mathbb{S}_2| = \pi_{\mathbb{S}_2}(n-1) \leq \sum_{i=0}^{d-1} \binom{n-1}{i}$.

Since $\binom{n-1}{d-1} + \binom{n-1}{d} = \binom{n}{d}$ and $\binom{n-1}{0} = \binom{n}{0}$

$$\begin{aligned}
\pi_{\mathbb{S}}(n) &\leq \pi_{\mathbb{S}_1}(n-1) + \pi_{\mathbb{S}_2}(n-1) \\
&\leq \binom{n-1}{0} + \binom{n-1}{1} + \cdots + \binom{n-1}{d} + \binom{n-1}{0} + \binom{n-1}{1} + \cdots + \binom{n-1}{d-1} \\
&\leq \binom{n-1}{0} + [\binom{n-1}{1} + \binom{n-1}{0}] + \cdots + [\binom{n-1}{d} + \binom{n-1}{d-1}] \\
&\leq \binom{n}{0} + \binom{n}{1} + \cdots + \binom{n}{d}.
\end{aligned}$$

■

6.6.4 Intersection Systems

Let $(U, \mathbb{S}_1)$ and $(U, \mathbb{S}_2)$ be two set systems on the same underlying set U . Define another set system, called the intersection system, $(U, \mathbb{S}_1 \cap \mathbb{S}_2)$, where $\mathbb{S}_1 \cap \mathbb{S}_2 = \{A \cap B \mid A \in \mathbb{S}_1; B \in \mathbb{S}_2\}$. In words, take the intersections of every set in $\mathbb{S}_1$ with every set in $\mathbb{S}_2$. A simple example is $U = \mathbb{R}^d$ and $\mathbb{S}_1$ and $\mathbb{S}_2$ are both the set of all half spaces. Then $\mathbb{S}_1 \cap \mathbb{S}_2$ consists of all sets defined by the intersection of two half spaces. This corresponds to taking the Boolean AND of the output of two threshold gates and is the most basic neural net besides a single gate. We can repeat this process and take the intersection of k half spaces. The following simple lemma helps us bound the growth of the shatter function as we do this.

Lemma 6.7 *Suppose $(U, \mathbb{S}_1)$ and $(U, \mathbb{S}_2)$ are two set systems on the same set U . Then*

$$\pi_{\mathbb{S}_1 \cap \mathbb{S}_2}(n) \leq \pi_{\mathbb{S}_1}(n) \pi_{\mathbb{S}_2}(n).$$

Proof: First observe that for $B \subseteq A$ if $A \cap S_1$ and $A \cap S_2$ are the same sets, then $B \cap S_1$ and $B \cap S_2$ must also be the same sets. Thus for $B \subseteq A$, $|\{B \cap S \mid S \in \mathbb{S}\}| \leq |\{A \cap S \mid S \in \mathbb{S}\}|$. The proof then follows from the fact that for any $A \subseteq U$, the number of sets of the form $A \cap (S_1 \cap S_2)$ with $S_1 \in \mathbb{S}_1$ and $S_2 \in \mathbb{S}_2$ is at most the number of sets of the form $A \cap S_1$ times the number of sets of the form $A \cap S_2$ since for fixed S_1 , $|(A \cap S_1) \cap S_2| \leq |A \cap S_2|$. ■

6.7 The VC Theorem

The VC theorem estimates the number of labeled training examples needed to train a good predictor of unlabeled test examples. Assume the examples are vectors and let U be the set of all vectors in the relevant space. Let $(U, \mathbb{H})$ be a set system. We assume that there is a subset $H \in \mathbb{H}$ according to which examples are labeled. Each $\mathbf{x} \in H$ gets the label +1 and each $\mathbf{x} \notin H$ gets the label -1. Our task is to learn a representation of H from a set of labeled training examples, so as to be able to predict labels of future examples. In learning theory, H is called a *concept* and in statistics it is called an *hypothesis*. Here, we will call it an hypothesis.

We are given a set S of points in U , each labeled according to an unknown hypothesis $H \in \mathbb{H}$. Our task is to learn H . In the case of half spaces, learning H exactly with a

finite number of training examples may not be possible. So we modify the task to learn a hypothesis H' which is approximately the same as H . How do we measure the difference between H and H' ? For half spaces, there is a natural notion of angle. For general $(U, \mathbb{H})$, there may not be such a notion. Even in the case of half spaces, if some regions of space are more important than others, angles and distances, which are the same everywhere in space, may not be the correct measure of approximation.

Valiant formulated the theoretical model of learning that gives an elegant answer to these issues. In Valiant's model, there is a probability distribution p , which the learning algorithm may not know. Training examples are picked in independent identical trials, each according to p . Each training example is labeled ± 1 as per an unknown hypothesis $H \in \mathbb{H}$. The purpose of learning is to come up with a hypothesis $H' \in \mathbb{H}$ that is used to predict labels of future test examples which are also picked in independent trials according to the same probability distribution p on U . The key insight here is to use the same probability to pick the training examples as the test examples.

Define the *prediction error* to be the probability that the label of a test example is predicted wrongly. Prediction error for a predictor H' is $p(H \triangle H')$ since the symmetric difference $H \triangle H'$ ⁹ is the set of examples on which the true hypothesis H and our predictor H' disagree. Since the learning algorithm only sees the set of training examples, the algorithm can come up with any hypothesis consistent with all the training examples. The central question is: How many training examples are sufficient so that any hypothesis consistent with all training examples makes prediction error of at most ε ? The training examples should rule out all possible H' with $p(H \triangle H') > \varepsilon$. For this, it is sufficient that at least one training example land in $H \triangle H'$ for every such H' . The H label of the example will be the opposite of its H' label ruling out H' .

The VC theorem below bounds the number of training examples needed in terms of the VC-dimension of the set system $(U, \mathbb{H})$. First a technical lemma needed in the proof.

Lemma 6.8 *If $y \geq x \ln x$, then*

$$\frac{2y}{\ln y} \geq x$$

provided $x > 4$.

Proof: First consider the situation where $y = x \ln x$. Then $\ln y = \ln x + \ln \ln x \leq 2 \ln x$. Thus $\frac{2y}{\ln y} \geq \frac{2x \ln x}{2 \ln x} \geq x$. It is easy to see by differentiation that $\frac{2y}{\ln y}$ is a monotonically increasing function of y in the range $y \geq x \ln x$ and $x > 4$. Thus, the lemma follows. ■

Theorem 6.9 (Vapnik-Chervonenkis Theorem) *Let $(U, \mathbb{H})$ be a set system with VC-dimension $d \geq 2$ and let p be a probability distribution on U . Let S be a set of m*

⁹For two sets H and H' , we denote their symmetric difference by $H \triangle H'$, namely, the set of elements belonging to precisely one of the sets H or H' .

independent samples picked from U according to p . For $m \geq \frac{1000d}{\varepsilon} \ln\left(\frac{d}{\varepsilon}\right)$, the probability that there exist sets H and H' in $\mathbb{H}$ with $p(H \triangle H') \geq \varepsilon$ for which $S \cap (H \triangle H') = \emptyset$ is at most $e^{-\varepsilon m/8}$.

Proof: The theorem asserts that a sufficiently large sample set S with high probability intersects every $H \triangle H'$ with $p(H \triangle H') \geq \varepsilon$. It is easy to prove that for a particular H and H' in $\mathbb{H}$, the probability that S misses $H \triangle H'$ is small. But $\mathbb{H}$ potentially has infinitely many sets. So, a union bound is not sufficient to prove the theorem.

We begin with an intuitive explanation of the proof. Let S_1 be a set of m samples and suppose for one pair H and H' in $\mathbb{H}$ with $p(H \triangle H') \geq \varepsilon$, the symmetric difference $H \triangle H'$ is missed by S_1 . According to p pick a second set S_2 of m samples independent of S_1 . With high probability, S_2 will have at least $\varepsilon m/2$ samples from the set $H \triangle H'$. Thus, if there is some H and H' in $\mathbb{H}$ with $p(H \triangle H') \geq \varepsilon$, which is missed by S_1 , then there is some H and H' in $\mathbb{H}$ with $p(H \triangle H') \geq \varepsilon$ whose symmetric difference is missed by S_1 , while S_2 has $\varepsilon m/2$ of its elements. The proof of the theorem lies in showing that this latter situation occurs with very low probability.

Instead of picking S_1 and then S_2 , pick a set W of $2m$ independent samples from U . Pick one of the $\binom{2m}{m}$ subsets of W of cardinality m uniformly at random for S_1 and let the remainder be S_2 . We claim that these two processes give the same distribution on S_1 and so we may use either process in our arguments. The proof uses each process at different places. Hence the technique is called double sampling.

Consider the second process of picking W and then S_1 and S_2 from W . Let H and H' be generic elements of $\mathbb{H}$. If $|W \cap (H \triangle H')| \geq \varepsilon m/2$, then it is highly unlikely that a random m -subset of W will completely miss $H \triangle H'$, so the probability of the event that $H \triangle H'$ is missed by S_1 , but $H \triangle H'$ has intersection at least $\varepsilon m/2$ with S_2 is very small. By a Chernoff bound, this probability falls off exponentially in m . Since this only works for a single pair of H and H' , We are still faced with the problem of the union bound over possibly infinitely many sets.

Once W is picked, we only need to worry about the $(H \triangle H') \cap W$ for all the H and H' in $\mathbb{H}$. Even though there may be infinitely many H and H' , the number of possible $H \cap W$ with $H \in \mathbb{H}$ is at most the shatter function of $2m$, which grows with $2m$ as a polynomial of degree equal to the VC dimension d of the set system. The number of possible $(H \triangle H') \cap W$ is at most the square of the number of possible $H \cap W$, since $(H \triangle H') \cap W = (H \cap W) \triangle (H' \cap W)$. So, we only need to ensure that the failure probability for each H and H' multiplied by a polynomial in $2m$ of degree $2d$ is $o(1)$. By a simple calculation $m \in \Omega((d/\varepsilon) \log(d/\varepsilon))$ suffices.

More formally, define two events E_1 and E_2 . Let E_1 be the event that there exist H

and H' and $T = H \triangle H'$, with $|T| \geq \varepsilon|U|$ and all points in S_1 miss T .

$$E_1 \quad \exists H \text{ and } H' \text{ in } \mathbb{H} \text{ with } |T| \geq \varepsilon|U| \text{ and } |T \cap S_1| = \emptyset$$

Let E_2 be the event that there exists an H and H' with $|T| \geq \varepsilon|U|$, all points in S_1 miss T , and S_2 intersects T in at least $\varepsilon m/2$ points. That is,

$$E_2 \quad \exists H \text{ and } H' \text{ in } \mathbb{H} \text{ with } |T| \geq \varepsilon|U|, |T \cap S_1| = \emptyset \text{ and } |T \cap S_2| \geq \frac{\varepsilon}{2}m.$$

We wish to show that $\text{Prob}(E_1)$ is very low. First we show that $\text{Prob}(E_2|E_1) \geq 1/2$. Then the bulk of the proof goes to show that $\text{Prob}(E_2)$ is very low. Since

$$\text{Prob}(E_2) \geq \text{Prob}(E_2|E_1)\text{Prob}(E_1)$$

this implies that $\text{Prob}(E_1)$ is very low.

We now show that $\text{Prob}(E_2|E_1) \geq 1/2$. Given E_1 there is a pair of sets H and H' such that $T \cap S_1 = \emptyset$. For this one pair H and H' , the probability that $S_2 \cap T \leq \varepsilon m/2$ is at most $1/2$ giving us $\text{Prob}(E_2|E_1) \geq 1/2$.

$\text{Prob}(E_2)$ is bounded by the double sampling technique. Instead of picking S_1 and then S_2 , pick a set W of $2m$ samples. Then pick a subset of size m out of W without replacement to be S_1 and let $S_2 = W \setminus S_1$. The distribution of S_1 and S_2 obtained this way is the same as picking S_1 and S_2 directly.

Now if E_2 occurs, then for some H and H' in $\mathbb{H}$, with $p(T) \geq \varepsilon$, we have both $|T \cap S_1| = 0$ and $|T \cap S_2| \geq \frac{\varepsilon}{2}m$. Since $|T \cap S_2| \geq \frac{\varepsilon}{2}m$ and $S_2 \subseteq W$, it follows that $|T \cap W| \geq \frac{\varepsilon}{2}m$. But if $|T \cap W| \geq \frac{\varepsilon}{2}m$ and S_1 is a random subset of cardinality m out of W , the probability that $|T \cap S_1| = 0$ is at most the probability of selecting m elements from the $2m - \frac{\varepsilon}{2}m$ elements other than the $\frac{\varepsilon}{2}m$ elements known to be in T . This probability is at most

$$\frac{\binom{2m - (\varepsilon/2)m}{m}}{\binom{2m}{m}} \leq \frac{m(m-1) \cdots (m - \frac{\varepsilon}{2}m + 1)}{(2m)(2m-1) \cdots (2m - \frac{\varepsilon}{2}m + 1)} \leq 2^{-\frac{\varepsilon m}{2}}.$$

This is the failure probability for just one pair H and H' . The number of possible $W \cap H$ is at most $\pi_{\mathbb{S}}(2m)$ which from Lemma 6.6 is at most $2(2m)^d \leq m^{2d}$ for $m \geq 4$. So the number of possible $(H \triangle H') \cap W$ is at most m^{4d} . By the union bound, the probability is at most $m^{4d} 2^{-\varepsilon m/2} \leq m^{4d} e^{-\varepsilon m/4}$. So we need to prove that $m^{4d} e^{-\varepsilon m/4} \leq e^{-\varepsilon m/8}$ or after some manipulation that $\frac{m}{\ln m} \geq \frac{32d}{\varepsilon}$. Apply Lemma 6.8 with $y = m$ and $x = 64d/\varepsilon$ to get the conclusion of the theorem. $\blacksquare$

6.8 Bibliographic Notes

Leslie Valiant formulated the theory of learning in a foundational paper- “A Theory of the learnable” [Val84]; this paper stipulates that the learning algorithm be measured on test

examples drawn from the same probability distribution from which the training examples were drawn. The connection between Valiant's learning model and the more classical notion of Vapnik-Chervonekis dimension [VC71] was struck by Blumer, Ehrenfrucht, Hausler and Warmuth [BEHW]. Boosting was first introduced by Schapire [Sch90]. A general reference on machine learning is the book [Mit97] and a more theoretical introduction is given in [KV95]. [SS01] is a reference on support vector machines and learning with kernels. The basic idea in boosting of taking a weighted majority of several decisions, where the weights get updated based on past experience has been used in economics as well as other areas. A survey of many applications of this general method can be found in [Aro11].

6.9 Exercises

Exercise 6.1 (Boolean OR has a linear separator) Take as examples all the 2^d elements of $\{0, 1\}^d$. Label the example by $+1$ if there is at least one coordinate with a $+1$ and label it by -1 if all its coordinates are 0 . This is like taking the Boolean OR, except that the coordinates are the real numbers 0 and 1 rather than true or false. Show that there is a linear separator for the labeled examples. Show that we can achieve a margin of $\Omega(1/d)$ for this problem.

Exercise 6.2 Repeat Exercise 6.1 for the AND function.

Exercise 6.3 Repeat Exercise 6.1 for majority and minority functions. [You may assume d is odd for this problem.]

Exercise 6.4 Show that the parity function, the Boolean function that is 1 if and only if an odd number of inputs is 1 , cannot be represented as a threshold function.

Exercise 6.5 Apply the perceptron learning algorithm to the following data.

$$\begin{array}{cccc} a_1 = [1, 2] & a_2 = [-2, -1] & a_3 = [-1, 1] & a_4 = [-1, -1] \\ l_1 = +1 & l_2 = -1 & l_3 = -1 & l_4 = +1 \end{array}$$

To simplify computation do not normalize the a_i to be unit vectors. Run the algorithm until $(w^T a_i)l_i > 0$ for all i . Plot the successive weight vectors without the threshold component.

Exercise 6.6 Suppose the starting $\mathbf{w}$ in the perceptron learning algorithm made an angle of 45° with the solution $\mathbf{w}^*$ whose margin is δ . Show that the number of iterations satisfies a smaller upper bound than $\frac{1}{\delta^2} - 1$ by a small modification to the proof of Theorem 6.1?

Exercise 6.7 The proof of Theorem 6.1 shows that for every $\mathbf{w}^*$, with $l_i(\mathbf{w}^{*T} \mathbf{a}_i) \geq \delta$ for $i = 1, 2, \dots, n$, the cosine of the angle between $\mathbf{w}$ and $\mathbf{w}^*$ is at least $\sqrt{t+1} \delta$ after t iterations. (So, the angle is at most $\cos^{-1}(\sqrt{t+1} \delta)$.) What happens if there are multiple $\mathbf{w}^*$, all satisfying $l_i(\mathbf{w}^{*T} \mathbf{a}_i) \geq \delta$ for $i = 1, 2, \dots, n$? Then, how can our one $\mathbf{w}$ make a small angle with all of these $\mathbf{w}^*$?

Exercise 6.8 Suppose the examples are points in d -space with $0, 1$ coordinates and the label of $\mathbf{x} \in \{0, 1\}^d$ is $+1$ if and only if $\mathbf{x} \neq \mathbf{0}$ and the least i for which $x_i = 1$ is odd. Otherwise the example's label is -1 . Show that the rule can be represented by the linear threshold function

$$(x_1, x_2, \dots, x_n) \left(1, -\frac{1}{2}, \frac{1}{4}, -\frac{1}{8}, \dots\right)^T = x_1 - \frac{1}{2}x_2 + \frac{1}{4}x_3 - \frac{1}{8}x_4 + \dots \geq 0$$

Exercise 6.9 (Hard) Prove that for the problem of Exercise 6.8, we cannot have a linear separator with margin at least $1/f(d)$ where $f(d)$ is bounded above by a polynomial function of d .

Exercise 6.10 (Hard) Recall the definition of margin where the linear separator is required to correctly classify all examples with a margin of at least δ . Suppose this condition is relaxed to say that the linear separator classifies all examples correctly and has a margin of at least δ on all but an ϵ fraction of the examples. Consider the following modified version of the perceptron learning algorithm:

Start with $\mathbf{w} = (1, 0, 0, \dots, 0)$

Repeat until $(\mathbf{w}\mathbf{a}_i)^T l_i > 0$ for all but at most 2ϵ fraction of the examples

Add to $\mathbf{w}$ the average of all $\mathbf{a}_i l_i$ with $(\mathbf{w}\mathbf{a}_i)^T l_i \leq 0$

Show that this is a “noise-tolerant” version of the algorithm. Namely, show that with the relaxed margin assumption, it correctly finds a linear separator that classifies all but at most a 2ϵ fraction correctly. Prove a bound on the number of steps the algorithm takes. What goes wrong if we use the old unmodified algorithm with the relaxed assumption?

Hint: Go over the proof of theorem 6.1 (convergence of perceptron learning algorithm) and adapt it. You need to modify the argument that the numerator increases in every step.

Exercise 6.11

1. Show that (6.3) can be reformulated as the unconstrained minimization of

$$|\mathbf{v}|^2 + c \sum_i \left(1 - l_i(\mathbf{v}\mathbf{a}_i)^T\right)^+.$$

2. Show that x^+ is a convex function.

The function x^+ does not have a derivative at 0. The function $(x^+)^2$ is smoother (its first derivative at 0 exists) and it is better to minimize

$$|\mathbf{v}|^2 + c \sum_i \left(\left(1 - l_i(\mathbf{v}\mathbf{a}_i)^T\right)^+\right)^2.$$

Exercise 6.12 Assume that the center of Figure 6.5 is $(0,0)$ and the side of each small square is of length 1. Show that a point has label +1 if and only if

$$(x_1 + 1)x_1(x_1 - 1)(x_2 + 1)x_2(x_2 - 1) \geq 0.$$

Consider only examples which are interior to a small square.

Exercise 6.13 Consider a set of examples in 2-dimensions where any example inside the circle $x_1^2 + x_2^2 = 1$ is labeled +1 and any example outside the circle is labeled -1. Construct a function φ so that the examples are linearly separable in the space $\varphi(\mathbf{x})$.

Exercise 6.14 Find a function φ that maps each of the regions below to a space where the regions are linearly separable.

1. $\{(x, y) \mid -1 \leq x \leq 1\}$
2. $\{(x, y) \mid -1 \leq x \leq 1, -1 \leq y \leq 1\}$

Exercise 6.15 (Hard) Label the points in the plane that are within the circle of radius one as $+1$. Label the points in the annulus of inner radius one and outer radius two as -1 and the points in an annulus of inner radius two and outer radius three as $+1$. Find a function φ mapping the points to a higher dimensional space where the two sets are linearly separable.

Exercise 6.16 Suppose examples are just real numbers in the interval $[0, 1]$ and suppose there are reals $0 < a_1 < a_2 < a_3 < \dots < a_k, 1$ an example is labeled $+1$ iff it is from $(0, a_1) \cup (a_2, a_3) \cup (a_4, a_5) \cup \dots$ [So alternate intervals are labeled $+1$.] Show that there is an embedding of the interval into an $O(k)$ dimensional space where we have a linear separator.

Exercise 6.17

1. Consider 2-dimensional points and let K be the kernel matrix where the entry for $\mathbf{a} = (a_1, a_2)$ and $\mathbf{b} = (b_1, b_2)$ is $(\mathbf{a}^T \mathbf{b})^2$. What is the mapping $\varphi(\mathbf{a})$ such that $K(\mathbf{a}, \mathbf{b}) = \varphi(\mathbf{a}) \cdot \varphi(\mathbf{b})$?
2. What mapping gives rise to the kernel $e^{-|\mathbf{a}-\mathbf{b}|^2}$?

Exercise 6.18 Let p be a polynomial of degree D with d variables. Prove the the number of monomials in the polynomial p is at most

$$\sum_{i=0}^D \binom{d+i-1}{d-1}.$$

Then prove that $\sum_{i=0}^D \binom{d+i-1}{d-1} \leq D(d+D)^{\min(d-1, D)}$.

Exercise 6.19 Produce a polynomial $p(x, y)$ whose arguments x and y are real numbers and a set of real numbers $a_1, a_2, \dots$ so that the matrix $K_{ij} = p(a_i, a_j)$ is not positive semidefinite.

Exercise 6.20 Are social networks sufficiently different so that a network can be identified by looking at a small region? One might try to answer this question by finding 100 regions in each of a number of social networks and training a support vector machine on 50 regions from each network labeled by the network they came from. The remaining 50 regions are then tested to see if the support vector machine can correctly classify most of them. If it can, then the networks must be different. To carry out this experiment, say that a region (subgraph) is formed by selecting a vertex at random along with all vertices within distance two of the selected vertex. The subgraph is then converted to a vector by defining a set of features such as the average degree of vertices in the subgraph. Create a list of 10 possible features.

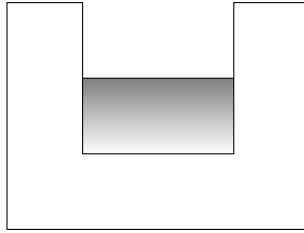


Figure 6.10: Data for exercise on boosting

Exercise 6.21 An $n \times n$ symmetric matrix A is positive semi-definite if for all x , $x^T A x \geq 0$. Prove that a positive semi-definite matrix A can be expressed as $A = BB^T$.

Exercise 6.22 Using boosting with threshold logic units find a solution to the data in Figure 6.10. The shaded area is data labeled $+1$ and the unshaded area is data labeled -1 .

Exercise 6.23 Make the proof that the majority of enough weak-learners in the boosting section is a strong learner rigorous by using inequalities instead of first order approximation. Prove that $T = \frac{3+\ln n}{\gamma \varepsilon} + \frac{1}{\varepsilon^2}$ will do for $\varepsilon < \gamma/8$

Exercise 6.24 (*Experts picking stocks*) Suppose there are n experts who are predicting whether one particular stock will go up or down at each of t time periods. There are only two outcomes at each time; up or down. You also have to make a prediction each time and after you do so, the actual outcome for that time period will be revealed to you. You may not assume any stochastic model of the experts (so past performance is no indication of the future). You are charged one for each wrong prediction. Can you pick nearly as well as the best expert where the best expert is the one who made the least number of wrong predictions overall? Show that boosting can help.

Exercise 6.25 Let $f(x) = ax^2 + bx + c$ and assume you measure $f(x)$ for a number of values of x . The measurements you make are corrupted by noise. You do not know f but would like to fit your data with some $g(x)$ so that you could predict $f(x)$ for additional values of x . If you use a high degree polynomial you will have more parameters than you need to fit the real data and will fit the noise. This is called over fitting. What should happen is as you go from degree 0 to 1 to 2 the predictive power should get better. As you increase the degree of g beyond 2 the predictive power should decrease. See if you can design an experiment to illustrate this effect of over fitting.

Exercise 6.26 What happens if in Section 6.5, instead of requiring

$$\text{Prob} \left(\left| |R \cap U| - \frac{n}{s} |R \cap S| \right| \leq \varepsilon n \text{ for every } R \right) \geq 1 - \delta,$$

one requires only:

$$\text{Prob} \left(\left| |R \cap U| - \frac{n}{s} |R \cap S| \right| \leq \varepsilon n \right) \geq 1 - \delta, \text{ for every } R ?$$

Exercise 6.27 Given n points in the plane and a circle C_1 containing at least three points (i.e., at least three points lie on or inside it) show that there exists a circle C_2 with two of the points on its circumference containing the same set of points as C_1 .

Exercise 6.28 Is the following statement true or false? Suppose we have n points in the plane and C_1 is a circle containing at least three points. There exists a circle C_2 with three points lying on the circle C_2 or two points lying on a diameter of C_2 and the set of points in C_2 is the same as the set of points in C_1 . Either give a counter example or a proof.

Exercise 6.29 Given n points in the plane define two circles as equivalent if they enclose the same set of points. Prove that there are only $O(n^3)$ equivalence classes of points defined by circles and thus only $O(n^3)$ subsets out of the 2^n subsets can be enclosed by circles.

Exercise 6.30 Prove that the VC-dimension of circles is three.

Exercise 6.31 Consider a 3-dimensional space.

1. What is the VC-dimension of rectangular boxes with axis-parallel sides?
2. What is the VC-dimension of d -dimensional rectangular boxes with axis-parallel sides?
3. What is the VC-dimension of spheres?

Exercise 6.32 Consider $d + 1$ points in general position that form a tetrahedron in d -space. Extend the $d - 1$ dimensional faces to hyper planes. the hyper planes partition d -space into a finite region and a number of infinite regions.

1. Select a point x in an infinite region that intersects the tetrahedron in a $d - 1$ dimensional face. Let v be the vertex not on the $d - 1$ dimensional selected face. Prove that the line from x to v intersects the $d - 1$ dimensional face.
2. Consider the 3-dimensional case. Select a point x in an infinite region that intersects the tetrahedron in an edge e . Let v_1 and v_2 be the two vertices not on the edge. Prove that the convex hull of $\{x, v_1, v_2\}$ intersects the edge. That is, one can draw a line from x to the edge (v_1, v_2) that intersects the edge e of the tetrahedron.
3. Prove that the convex hull of a point in an infinite region plus the vertices not on the face of the tetrahedron that meets the infinite region intersects the convex hull of the face.

Exercise 6.33 (Squares) Show that there is a set of three points which can be shattered by axis-parallel squares. Show that the system of axis-parallel squares cannot shatter any set of four points.

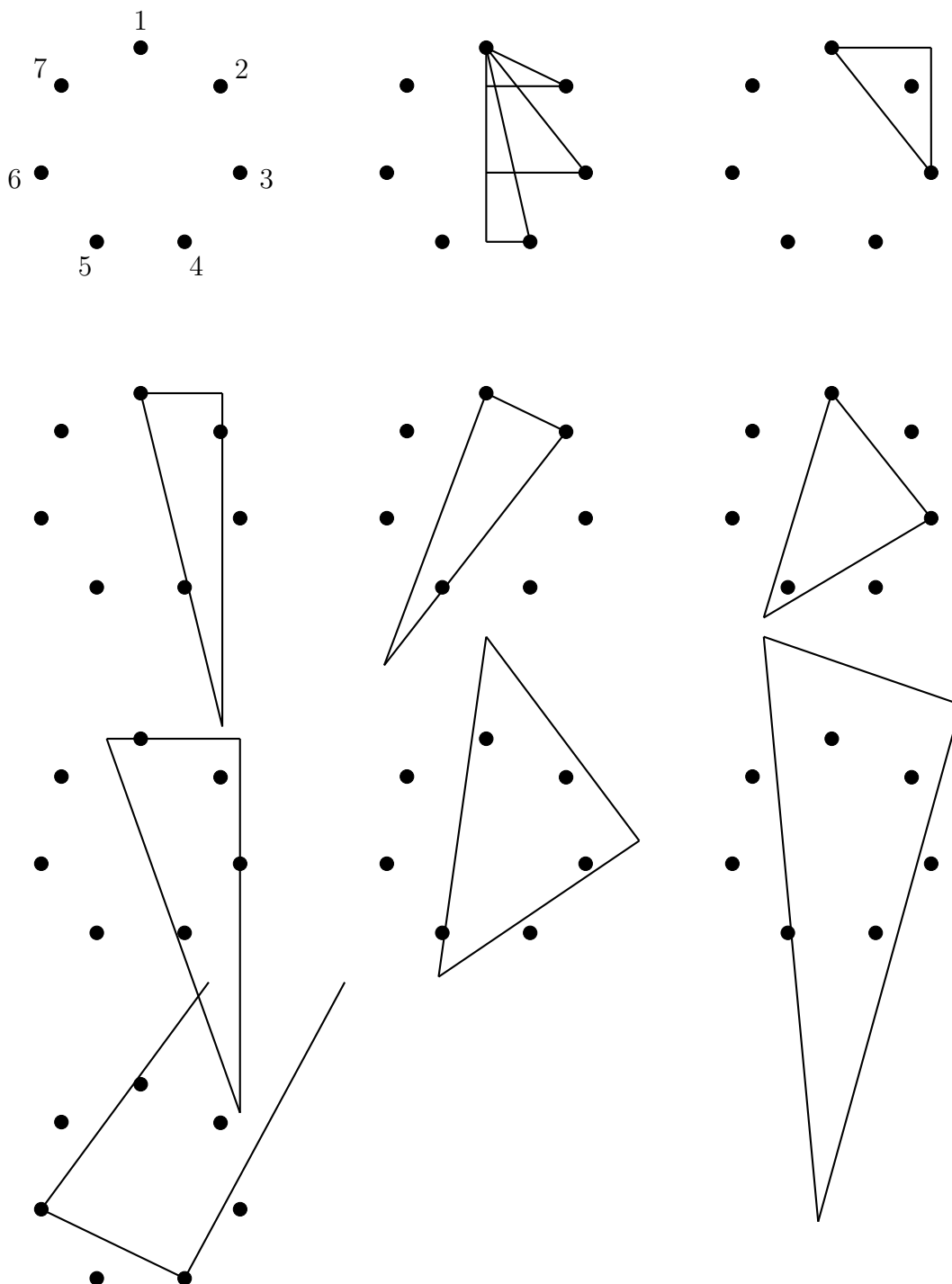
Exercise 6.34 *Show that the VC-dimension of axis-aligned right triangles with the right angle in the lower left corner is four.*

Exercise 6.35 *Prove that the VC-dimension of 45° , 45° , 90° triangles with right angle in the lower left is four.*

Exercise 6.36 *Show that the VC-dimension of arbitrary right triangles is seven.*

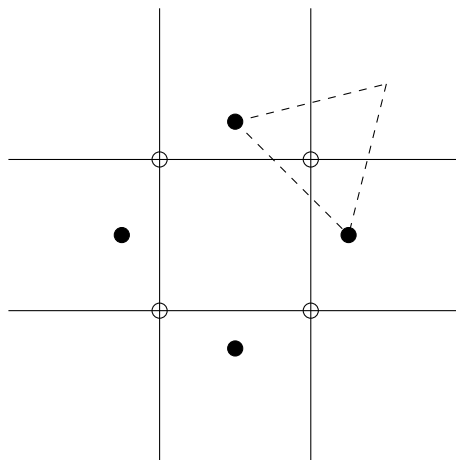
Consider a set of seven points on a circle. Any subset of zero, one, six, or seven points can obviously be separated. For all other subsets assume vertex one is in the set and is the counterclockwise most vertex of the largest contiguous subset by a symmetry argument in which the set is rotated and possibly flipped.

1100000	1010000	1001000	
1110000	1101000	1100100	1010100
1111000	1110100	1101100	1101010



If a set of eight points can be shattered no point can lie in the convex hull of the other seven. Hence they must lie on the boundary of a convex shape. Consider four alternate points as one walks around the boundary. These cannot be separate. To see this draw a quadrilateral with these four points as boundary. Extend the edges to infinity dividing the plane into a finite region and eight infinite regions. Draw a triangle containing the four other points. No vertex of the triangle can lie in an infinite region that meets the

quadrilateral in a vertex otherwise that vertex would be inside the triangle. There must be a vertex in each region that meets the quadrilateral in an edge. But there are four such regions and hence a contradiction.



Exercise 6.37 *What is the VC-dimension of triangles?*

Exercise 6.38 *Prove that the VC dimension of convex polygons is infinite.*

Exercise 6.39 *If a class contains only convex sets prove that it cannot shatter any set in which some point is in the convex hull of other points in the set.*

Exercise 6.40

1. *Prove that no set of six points can be shattered by squares in arbitrary position (rotation allowed).*
2. *Show that the VC-dimension of squares in arbitrary position is 5.*

Exercise 6.41

1. *Show that the set of seven vertices of a regular heptagon can be shattered by rotated rectangles.*
2. *Prove that no set of eight points can be shattered by rotated rectangles there by showing that the VC-dimension of rectangles in arbitrary position is 7.*

Exercise 6.42 *What is the VC-dimension of the family of quadrants? A quadrant Q is a set of points of one of the four types below:*

1. $Q = \{(x, y) : (x - x_0, y - y_0) \geq (0, 0)\},$
2. $Q = \{(x, y) : (x_0 - x, y - y_0) \geq (0, 0)\},$

3. $Q = \{(x, y) : (x_0 - x, y_0 - y) \geq (0, 0)\}$, or

4. $Q = \{(x, y) : (x - x_0, y_0 - y) \geq (0, 0)\}$.

Exercise 6.43 Create a list of simple shapes for which we can calculate the VC-dimension and indicate the VC-dimension for each shape on your list.

Exercise 6.44 For large n , how should you place n points on the plane so that the maximum number of subsets of the n points are defined by rectangles? Can you achieve $4n$ subsets of size 2? Can you do better? What about size 3? What about size 10?

Exercise 6.45 For large n , how should you place n points on the plane so that the maximum number of subsets of the n points are defined by

1. half spaces?
2. circles?
3. axis-parallel rectangles?
4. some other simple shape of your choosing?

For each of the shapes how many subsets of size two, three, etc can you achieve?

Exercise 6.46 What is the shatter function for 2-dimensional half spaces? That is, given n points in the plane, how many subsets can be defined by half spaces?

Exercise 6.47 What does it mean to shatter the empty set? How many subsets does one get?

Exercise 6.48 Intuitively define the most general form of a set system of VC-dimension one. Give an example of such a set system that can generate n subsets of an n element set. What is the form of the most general set system of dimension two.

Exercise 6.49 (Hard) We proved that if the VC-dimension is small, then the shatter function is small as well. Can you prove some sort of converse to this?

Exercise 6.50 If $(U, \mathbb{S}_1), (U, \mathbb{S}_2), \dots, (U, \mathbb{S}_k)$ are k set systems on the same ground set U show that $\pi_{\mathbb{S}_1 \cap \mathbb{S}_2 \cap \dots \mathbb{S}_k}(n) \leq \pi_{\mathbb{S}_1}(n) \pi_{\mathbb{S}_2}(n) \cdots \pi_{\mathbb{S}_k}(n)$.

Exercise 6.51 Show that in the “double sampling” procedure, the probability of picking a pair of multi-sets T and T' , each of cardinality m , by first picking T and then T' is the same as picking a W of cardinality $2m$ and then picking uniformly at random a subset T out of W of cardinality m and letting T' be $W - T$. For this exercise, assume that p , the underlying probability distribution is discrete.

Exercise 6.52 Randomly select n integers from the set $\{1, 2, \dots, 2n\}$ without replacement. In the limit as n goes to infinity, what is the probability of not selecting any integer in the set $\{1, 2, \dots, k\}$ for k a constant independent of n ? For $k = \ln n$?

Exercise 6.53 Write a short paragraph about what you learned in this chapter.

7 Algorithms for Massive Data Problems

Massive Data, Sampling

This chapter deals with massive data problems where the input data, a graph, a matrix or some other object, is too large to be stored in random access memory. One model for such problems is the streaming model, where the data can be seen only once. In the streaming model, the natural technique to deal with the massive data is sampling. Sampling is done “on the fly”. As each piece of data is seen, based on a coin toss, one decides whether to include the data in the sample. Typically, the probability of including the data point in the sample may depend on its value. Models allowing multiple passes through the data are also useful; but the number of passes needs to be small. We always assume that random access memory, RAM, is limited, so the entire data cannot be stored in RAM.

To introduce the basic flavor of sampling on the fly, consider the following primitive. From a stream of n positive real numbers, $a_1, a_2, \dots, a_n$, draw a sample element a_i so that the probability of picking an element is proportional to its value. It is easy to see that the following sampling method works. Upon seeing $a_1, a_2, \dots, a_i$, keep track of the sum $a = a_1 + a_2 + \dots + a_i$ and a sample a_j , $j \leq i$, drawn with probability proportional to its value. On seeing a_{i+1} , replace the current sample by a_{i+1} with probability $\frac{a_{i+1}}{a + a_{i+1}}$ and update a . If the current element is replaced by a_{i+1} , then clearly a_{i+1} is selected with probability proportional to its value. If the current element is not replaced, then a_j is the selected element and its probability of having been selected is

$$\frac{a_j}{a_1 + a_2 + \dots + a_i} \left(1 - \frac{a_{i+1}}{a_1 + a_2 + \dots + a_{i+1}} \right) = \frac{a_j}{a_1 + a_2 + \dots + a_{i+1}}.$$

7.1 Frequency Moments of Data Streams

An important class of problems concerns the frequency moments of data streams. Here a data stream $a_1, a_2, \dots, a_n$ of length n consists of symbols a_i from an alphabet of m possible symbols, which for convenience we denote as $\{1, 2, \dots, m\}$. Throughout this section, n, m , and a_i will have these meanings and s , for symbol, will denote a generic element of $\{1, 2, \dots, m\}$. The frequency f_s of the symbol s is the number of occurrences of s in the stream. For a nonnegative integer p , the p^{th} frequency moment of the stream is

$$\sum_{s=1}^m f_s^p.$$

Note that the $p = 0$ frequency moment corresponds to the number of distinct symbols occurring in the stream. The first frequency moment is just n , the length of the string. The second frequency moment, $\sum_s f_s^2$, is useful in computing the variance of the stream.

$$\frac{1}{m} \sum_{s=1}^m \left(f_s - \frac{n}{m} \right)^2 = \frac{1}{m} \sum_{s=1}^m \left(f_s^2 - 2 \frac{n}{m} f_s + \left(\frac{n}{m} \right)^2 \right) = \frac{1}{m} \sum_{s=1}^m f_s^2 - \frac{n^2}{m^2}$$

In the limit as p becomes large, $\left(\sum_{s=1}^m f_s^p\right)^{1/p}$ is the frequency of the most frequent element(s).

We will describe sampling based algorithms to compute these quantities for streaming data shortly. But first a note on the motivation for these various problems. The identity and frequency of the the most frequent item or more generally, items whose frequency exceeds a fraction of n , is clearly important in many applications. If the items are packets on a network with source destination addresses, the high frequency items identify the heavy bandwidth users. If the data is purchase records in a supermarket, the high frequency items are the best-selling items. Determining the number of distinct symbols is the abstract version of determining such things as the number of accounts, web users, or credit card holders. The second moment and variance are useful in networking as well as in database and other applications. Large amounts of network log data are generated by routers that can record for all the messages passing through them, the source address, destination address, and the number of packets. This massive data cannot be easily sorted or aggregated into totals for each source/destination. But it is important to know if a few popular source-destination pairs generate a lot of the traffic for which the second moment is the natural measure.

7.1.1 Number of Distinct Elements in a Data Stream

Consider a sequence $a_1, a_2, \dots, a_n$ of n elements, each a_i an integer in the range 1 to m where n and m are very large. Suppose we wish to determine the number of distinct a_i in the sequence. Each a_i might represent a credit card number extracted from a sequence of credit card transactions and we wish to determine how many distinct credit card accounts there are. The model is a data stream where symbols are seen one at a time. We first show that any deterministic algorithm that determines the number of distinct elements exactly must use at least m bits of memory.

Lower bound on memory for exact deterministic algorithm

Suppose we have seen the first $k \geq m$ symbols of the stream. The set of distinct symbols seen so far could be any of the 2^m subsets of $\{1, 2, \dots, m\}$. Each subset must result in a different state for our algorithm and hence m bits of memory are required. To see this, suppose first that two different size subsets of distinct symbols lead to the same internal state. Then our algorithm would produce the same count of distinct symbols for both inputs, clearly an error for one of the input sequences. If two sequences with the same number of distinct elements but different subsets lead to the same state, then seeing a symbol that appeared in one sequence but not the other would result in subsets of different sizes and thus require different states.

Algorithm for the Number of distinct elements

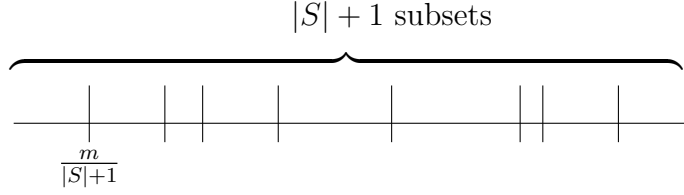


Figure 7.1: Estimating the size of S from the minimum element in S which has value approximately $\frac{m}{|S|+1}$. The elements of S partition the set $\{1, 2, \dots, m\}$ into $|S| + 1$ subsets each of size approximately $\frac{m}{|S|+1}$.

In estimating a quantity such as the number of distinct elements, we are interested in the relative error not the absolute error. For example, we may want to estimate a quantity so that our estimate is within a multiplicative factor of $1 \pm \epsilon$ of the correct value. Thus, the variance of our estimate may be as large as the square of the quantity itself. We can reduce the variance by multiple independent estimates. With $\frac{1}{\epsilon^2}$ independent estimates, the variance is reduced by a multiplicative factor of ϵ^2 and the standard deviation is reduced by a multiplicative factor of ϵ . Thus with $\frac{1}{\epsilon}$ independent estimates, the estimate is highly likely to be within a multiplicative constant of $1 + c\epsilon$ of the correct value for some small constant c .

Let $a_1, a_2, \dots, a_n$ be a sequence of elements where each $a_i \in \{1, 2, \dots, m\}$. The number of distinct elements can be estimated with $O(\log m)$ space. Let $S \subseteq \{1, 2, \dots, m\}$ be the set of elements that appear in the sequence. Suppose that the elements of S were selected uniformly at random from $\{1, 2, \dots, m\}$. Let $\min$ denote the minimum element of S . Knowing the minimum element of S allows us to estimate the size of S . The elements of S partition the set $\{1, 2, \dots, m\}$ into $|S| + 1$ subsets each of size approximately $\frac{m}{|S|+1}$. See Figure 7.1. Thus, the minimum element of S should have value close to $\frac{m}{|S|+1}$. Solving $\min = \frac{m}{|S|+1}$ yields $|S| = \frac{m}{\min} - 1$. Since we can determine $\min$, this gives us an estimate of $|S|$.

The above analysis required that the elements of S were picked uniformly at random from $\{1, 2, \dots, m\}$. This is generally not the case when we have a sequence $a_1, a_2, \dots, a_n$ of elements from $\{1, 2, \dots, m\}$. Clearly if the elements of S were obtained by selecting the $|S|$ smallest elements of $\{1, 2, \dots, m\}$, the above technique would give the wrong answer. If the elements are not picked uniformly at random, can we estimate the number of distinct elements? The way to solve this problem is to use a hash function h where

$$h : \{1, 2, \dots, m\} \rightarrow \{0, 1, 2, \dots, M - 1\}$$

To count the number of distinct elements in the input, count the number of elements in the mapped set $\{h(a_1), h(a_2), \dots\}$. The point being that $\{h(a_1), h(a_2), \dots\}$ behaves

like a random subset and so the above heuristic argument using the minimum to estimate the number of elements may apply. If we needed $h(a_1), h(a_2), \dots$ to be completely independent, the space needed to store the hash function would be too high. Fortunately, only 2-way independence is needed. We recall the formal definition of 2-way independence below. But first recall that a hash function is always chosen at random from a family of hash functions and phrases like “probability of collision” refer to the probability in the choice of hash function.

Universal Hash Functions

A set of hash functions

$$H = \{h \mid h : \{1, 2, \dots, m\} \rightarrow \{0, 1, 2, \dots, M-1\}\}$$

is *2-universal* if for all x and y in $\{1, 2, \dots, m\}$, $x \neq y$, and for all z and w in $\{0, 1, 2, \dots, M-1\}$

$$\text{Prob}(h(x) = z \text{ and } h(y) = w) = \frac{1}{M^2}$$

for a randomly chosen h . The concept of a 2-universal family of hash functions is that given x , $h(x)$ is equally likely to be any element of $\{0, 1, 2, \dots, M-1\}$ and for $x \neq y$, $h(x)$ and $h(y)$ are independent.

We now give an example of a 2-universal family of hash functions. For simplicity let M be a prime. For each pair of integers a and b in the range $[0, M-1]$, define a hash function

$$h_{ab}(x) = ax + b \pmod{M}$$

To store the hash function h_{ab} , store the two integers a and b . This requires only $O(\log M)$ space. To see that the family is 2-universal note that $h(x) = z$ and $h(y) = w$ if and only if

$$\begin{pmatrix} x & 1 \\ y & 1 \end{pmatrix} \begin{pmatrix} a \\ b \end{pmatrix} = \begin{pmatrix} z \\ w \end{pmatrix} \pmod{M}$$

If $x \neq y$, the matrix $\begin{pmatrix} x & 1 \\ y & 1 \end{pmatrix}$ is invertible modulo M and there is only one solution for a and b . Hence, for a and b chosen uniformly at random, the probability of the equation holding is exactly $\frac{1}{M^2}$. Thus,

$$\text{Prob}(h_{ab}(x) = z \text{ and } h_{ab}(y) = w) = \text{Prob}(h_{ab}(x) = z) \text{Prob}(h_{ab}(y) = w)$$

and $h_{ab}(x)$ and $h_{ab}(y)$ are statistically independent.

Analysis of distinct element counting algorithm

Let $b_1, b_2, \dots, b_d$ be the distinct values that appear in the input. Then $S = \{h(b_1), h(b_2), \dots, h(b_d)\}$ is a set of d random and 2-way independent values from the set $\{0, 1, 2, \dots, M-1\}$. We now show that $\frac{M}{\min}$ is a good estimate for d , the number of distinct elements in the input, where $\min$ is the minimum value in the set S .

Lemma 7.1 Assume $M > 100d$. With probability at least $\frac{2}{3}$, $\frac{d}{6} \leq \frac{M}{\min} \leq 6d$, where $\min$ is the smallest element of S .

Proof: First, we show that $\text{Prob} \left[\frac{M}{\min} > 6d \right] < \frac{1}{6}$.

$$\text{Prob} \left[\frac{M}{\min} > 6d \right] = \text{Prob} \left[\min < \frac{M}{6d} \right] = \text{Prob} \left[\exists k, h(b_k) < \frac{M}{6d} \right]$$

For $i = 1, 2, \dots, d$, define the indicator variable

$$z_i = \begin{cases} 1 & \text{if } h(b_i) < \frac{M}{6d} \\ 0 & \text{otherwise} \end{cases}$$

and let $z = \sum_{i=1}^d z_i$. If $h(b_i)$ is chosen randomly from $\{0, 1, 2, \dots, M-1\}$, then $\text{Prob}(z_i = 1) < \frac{1}{6d}$. Thus, $E(z_i) < \frac{1}{6d}$ and $E(z) < \frac{1}{6}$. Now

$$\begin{aligned} \text{Prob} \left[\frac{M}{\min} > 6d \right] &= \text{Prob} \left[\min < \frac{M}{6d} \right] \\ &= \text{Prob} \left[\exists k, h(b_k) < \frac{M}{6d} \right] \\ &\leq \text{Prob}(z \geq 1) \\ &\leq \text{Prob}(z \geq 6E(z)). \end{aligned}$$

By Markov's inequality $\text{Prob}(z \geq 6E(z)) \leq \frac{1}{6}$.

Finally, we show that $\text{Prob} \left(\frac{M}{\min} < \frac{d}{6} \right) < \frac{1}{6}$.

$$\text{Prob} \left[\frac{M}{\min} < \frac{d}{6} \right] = \text{Prob} \left(\min > \frac{6M}{d} \right) = \text{Prob} \left(\forall k, h(b_k) > \frac{6M}{d} \right)$$

For $i = 1, 2, \dots, d$ define the indicator variable

$$y_i = \begin{cases} 0 & \text{if } h(b_i) > \frac{6M}{d} \\ 1 & \text{otherwise} \end{cases}$$

and let $y = \sum_{i=1}^d y_i$. Now $\text{Prob}(y_i = 1) > \frac{6}{d}$, $E(y_i) > \frac{6}{d}$, and $E(y) > 6$. For 2-way independent random variables, the variance of their sum is the sum of their variances. So $\text{Var}(y) = d\text{Var}(y_1)$. Further, since y_1 is 0 or 1

$$\text{Var}(y_1) = E[(y_1 - E(y_1))^2] = E(y_1^2) - E^2(y_1) = E(y_1) - E^2(y_1) \leq E(y_1).$$

Thus, $\text{Var}(y) \leq E(y)$. Now by the Chebychev inequality,

$$\begin{aligned} \text{Prob} \left[\frac{M}{\min} < \frac{d}{6} \right] &= \text{Prob} \left[\min > \frac{6M}{d} \right] = \text{Prob} \left[\forall k \ h(b_k) > \frac{6M}{d} \right] \\ &= \text{Prob}(y = 0) \leq \text{Prob}[|y - E(y)| \geq E(y)] \\ &\leq \frac{\text{Var}(y)}{E^2(y)} \leq \frac{1}{E(y)} \leq \frac{1}{6}. \end{aligned}$$

Since $\frac{M}{\min} > 6d$ with probability at most $\frac{1}{6}$ and $\frac{M}{\min} < \frac{d}{6}$ with probability at most $\frac{1}{6}$, $\frac{d}{6} \leq \frac{M}{\min} \leq 6d$ with probability at least $\frac{2}{3}$. ■

7.1.2 Counting the Number of Occurrences of a Given Element.

To count the number of occurrences of an element in a stream requires at most $\log n$ space where n is the length of the stream. Clearly, for any length stream that occurs in practice, we can afford $\log n$ space. For this reason, the following material may never be used in practice, but the technique is interesting and may give insight into how to solve some other problems.

Consider a string of 0's and 1's of length n in which we wish to count the number of occurrences of 1's. Clearly if we had $\log n$ bits of memory we could keep track of the exact number of 1's. However, we can approximate the number with only $\log \log n$ bits.

Let m be the number of 1's that occur in the sequence. Keep a value k such that 2^k is approximately the number of occurrences m . Storing k requires only $\log \log n$ bits of memory. The algorithm works as follows. Start with $k=0$. For each occurrence of a 1, add one to k with probability $1/2^k$. At the end of the string, the quantity $2^k - 1$ is the estimate of m . To obtain a coin that comes down heads with probability $1/2^k$, flip a fair coin, one that comes down heads with probability $1/2$, k times and report heads if the fair coin comes down heads in all k flips.

Given k , on average it will take 2^k ones before k is incremented. Thus, the expected number of 1's to produce the current value of k is $1 + 2 + 4 + \dots + 2^{k-1} = 2^k - 1$.

7.1.3 Counting Frequent Elements

The Majority and Frequent Algorithms

First consider the very simple problem of n people voting where there are m candidates, $\{1, 2, \dots, m\}$. We want to determine if one candidate gets a majority vote and if so who. Formally, we are given a stream of integers $a_1, a_2, \dots, a_n$, each a_i belonging to $\{1, 2, \dots, m\}$. We want to determine whether there is some s in $\{1, 2, \dots, m\}$ which occurs more than $n/2$ times and if so which s . To solve the problem exactly with a deterministic algorithm on streaming read only once data where $m > n$, requires $\Omega(n)$ space. Suppose

n is even and the first $n/2$ items are all distinct and the last $n/2$ items are identical. After reading the first $n/2$ items, we need to remember exactly which elements of $\{1, 2, \dots, m\}$ have occurred. If for two different sets of elements occurring in the first half of the stream, the contents of the memory are the same, then a mistake would occur if the second half of the stream consists solely of an element in one set, but not the other. Thus, $\log_2 \binom{m}{n/2}$ bits of memory, which if $m > n$ is $\Omega(n)$, are needed.

The following is a simple low-space algorithm that always finds the majority vote if there is one. If there is no majority vote, the output may be arbitrary. That is, there may be “false positives”, but no “false negatives”.

Majority Algorithm

Store a_1 and initialize a counter to one. For each subsequent a_i , if a_i is the same as the currently stored item, increment the counter by one. If it differs, decrement the counter by one provided the counter is nonzero. If the counter is zero, then store a_i and set the counter to one.

If s is not the stored element at the end, each copy of s either incremented or decremented the counter. Associate with each copy of s another symbol as follows. If s incremented the counter, associate with s the symbol that deleted that increment from the counter. If s decremented the counter, then associate with s the symbol responsible for that element of the count. If s is not the stored element at the end, then each copy of s is associated with another symbol in the sequence and thus the number of copies of s is at most $\frac{n}{2}$. Thus, if there is a majority element, it must be stored on the counter at the end.

Next we modify the above algorithm so that not just the majority, but also items with frequency above some threshold are detected. We will also ensure (approximately) that there are no false positives as well as no false negatives. Indeed the algorithm below will find the frequency (number of occurrences) of each element of $\{1, 2, \dots, m\}$ to within an additive term of $\frac{n}{k+1}$ using $O(k \log n)$ space by keeping k counters instead of just one counter.

Algorithm Frequent

Maintain a list of items being counted. Initially the list is empty. For each item, if it is the same as some item on the list, increment its counter by one. If it differs from all the items on the list, then if there are less than k items on the list, add the item to the list with its counter set to one. If there are already k items on the list decrement each of the current counters by one. Delete an element from the list if its count becomes zero.

Theorem 7.2 *At the end of Algorithm Frequent, for each s in $\{1, 2, \dots, m\}$, its counter on the list is at least the number of occurrences of s in the stream minus $n/(k+1)$. In*

particular, if some s does not occur on the list, its counter is zero and the theorem asserts that it occurs fewer than $n/(k+1)$ times in the stream.

Proof: View each decrement counter step as eliminating some items. An item is eliminated if it is the current a_i being read and there are already k symbols different from it on the list in which case it and k other items are simultaneously eliminated. At the end of the stream any element not eliminated contributes to one of the counts, The elimination of each occurrence of an s in $\{1, 2, \dots, m\}$ is really the elimination of $k+1$ items. Thus, no more than $n/(k+1)$ occurrences of any symbol can be eliminated. If an item is not eliminated, then it must still be on the list at the end. This proves the theorem. ■

Theorem 7.2 implies that we can compute the true relative frequency, the number of occurrences divided by n , of every s in $\{1, 2, \dots, m\}$ to within an additive term of $\frac{n}{k+1}$.

7.1.4 The Second Moment

This section focuses on computing the second moment of a stream with symbols from $\{1, 2, \dots, m\}$. Let f_s denote the number of occurrences of symbol s in the stream. The second moment of the stream is given by $\sum_{s=1}^m f_s^2 = f_1^2 + f_2^2 + \dots + f_m^2$. One could compute each sum separately but this would require m counters and we would like to compute $\sum_{s=1}^m f_s^2$ with $\log m$ space. Instead to calculate the second moment, for each symbol s , $1 \leq s \leq m$, independently set a random variable x_s to ± 1 with probability $1/2$. Maintain a single sum by adding x_s to the sum each time the symbol s occurs in the stream. At the end of the stream, the sum will equal $\sum_{s=1}^m x_s f_s$. The expected value of the sum will be zero where the expectation is over the choice of the ± 1 values for the x_s .

$$E \left(\sum_{s=1}^m x_s f_s \right) = 0.$$

Although the expected value of the sum is zero, its actual value is a random variable and the expected value of the square of the sum is given by

$$E \left(\sum_{s=1}^m x_s f_s \right)^2 = E \left(\sum_{s=1}^m x_s^2 f_s^2 \right) + 2E \left(\sum_{s \neq t} x_s x_t f_s f_t \right) = \sum_{s=1}^m f_s^2,$$

The last equality follows since $E(x_s x_t) = 0$ for $s \neq t$. Thus

$$a = \left(\sum_{s=1}^m x_s f_s \right)^2$$

is an estimator of $\sum_{s=1}^m f_s^2$. One difficulty which we will come back to is that to store all the x_i requires space m and we want to do the calculation in $\log m$ space.

How good this estimator is depends on its variance which we now compute.

$$\text{Var}(a) \leq E \left(\sum_{s=1}^m x_s f_s \right)^4 = E \left(\sum_{1 \leq s, t, u, v \leq m} x_s x_t x_u x_v f_s f_t f_u f_v \right)$$

The first inequality is because the variance is at most the second moment and the second equality is by expansion. In the second sum, since the x_s are independent, if any one of s, u, t , or v is distinct from the others, then the expectation of the whole term is zero. Thus, we need to deal only with terms of the form $x_s^2 x_t^2$ for $t \neq s$ and terms of the form x_s^4 . Note that this does not need the full power of mutual independence of all the x_s , it only needs 4-way independence, that any four of the x'_s are mutually independent. In the above sum, there are four indices s, t, u, v and there are $\binom{4}{2}$ ways of choosing two of them that have the same x value. Thus,

$$\begin{aligned} \text{Var}(a) &\leq \binom{4}{2} E \left(\sum_{s=1}^m \sum_{t=s+1}^m x_s^2 x_t^2 f_s^2 f_t^2 \right) + E \left(\sum_{s=1}^m x_s^4 f_s^4 \right) \\ &= 6 \sum_{s=1}^m \sum_{t=s+1}^m f_s^2 f_t^2 + \sum_{s=1}^m f_s^4 \\ &\leq 3 \left(\sum_{s=1}^m f_s^2 \right)^2. \end{aligned}$$

The variance can be reduced by a factor of r , which reduces the standard deviation by $\sqrt{r}$, by taking the average of r independent trials. With r independent trials the variance would be at most $\frac{3}{r} E^2(a)$, so to achieve relative error ε in the estimate of $\sum_{s=1}^m f_s^2$, $O(1/\varepsilon^2)$ independent trials suffice.

We briefly discuss the independent trials here, so as to understand exactly the amount of independence needed. Instead of computing a using the running sum $\sum_{s=1}^m x_s f_s$ for one random vector $\mathbf{x} = (x_1, x_2, \dots, x_m)$, independently generate r m -vectors $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(r)}$ at the outset and compute r running sums

$$\sum_{s=1}^m x_s^{(1)} f_s, \sum_{s=1}^m x_s^{(2)} f_s, \dots, \sum_{s=1}^m x_s^{(r)} f_s.$$

Let $a_1 = \left(\sum_{s=1}^m x_s^{(1)} f_s \right)^2$, $a_2 = \left(\sum_{s=1}^m x_s^{(2)} f_s \right)^2$, $\dots$, $a_r = \left(\sum_{s=1}^m x_s^{(r)} f_s \right)^2$. Our estimate is $\frac{1}{r}(a_1 + a_2 + \dots + a_r)$. The variance of this estimator is

$$\text{Var} \left[\frac{1}{r} (a_1 + a_2 + \dots + a_r) \right] = \frac{1}{r^2} [\text{Var}(a_1) + \text{Var}(a_2) + \dots + \text{Var}(a_r)] = \frac{1}{r} \text{Var}(a_1),$$

where we have assumed that the $a_1, a_2, \dots, a_r$ are mutually independent. Now we compute the variance of a_1 as we have done for the variance of a . Note that this calculation assumes only 4-way independence between the coordinates of $\mathbf{x}^{(1)}$. We summarize the assumptions here for future reference:

To get an estimate of $\sum_{s=1}^m f_s^2$ within relative error ε with probability close to one, say at least 0.9999, it suffices to have $r = O(1/\varepsilon^2)$ vectors $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(r)}$, each with m coordinates of ± 1 with

1. $E(\mathbf{x}^{(1)}) = E(\mathbf{x}^{(2)}) = \dots = E(\mathbf{x}^{(r)}) = \mathbf{0}$.
2. $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(r)}$ are mutually independent. That is, for any r vectors $\mathbf{v}^{(1)}, \mathbf{v}^{(2)}, \dots, \mathbf{v}^{(r)}$ with ± 1 coordinates, $\text{Prob}(\mathbf{x}^{(1)} = \mathbf{v}^{(1)}, \mathbf{x}^{(2)} = \mathbf{v}^{(2)}, \dots, \mathbf{x}^{(r)} = \mathbf{v}^{(r)}) = \frac{1}{2^{mr}}$.
3. Any four coordinates of $\mathbf{x}^{(1)}$ are independent. I.e., for any distinct s, t, u , and v in $\{1, 2, \dots, m\}$ and any a, b, c , and d in $\{-1, +1\}$,

$$\text{Prob}\left(x_s^{(1)} = a, x_t^{(1)} = b, x_u^{(1)} = c, x_v^{(1)} = d\right) = \frac{1}{16}.$$

Same for $\mathbf{x}^{(2)}, \mathbf{x}^{(3)}, \dots, \mathbf{x}^{(r)}$.

In fact, (1) follows from (3). The reader can prove this.

The only drawback with the algorithm described here is the need to keep the r vectors $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(r)}$ in memory in order to do the running sums. This is too space-expensive. We need to solve the problem in space dependent upon the logarithm of the size of the alphabet m , not m itself. If ε is $\Omega(1)$, then r is $O(1)$, so it is not the number of trials r which is the problem. It is the m .

In the next section, we will see that the computation can be done in $O(\log m)$ space by using pseudo-random vectors $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(r)}$ instead of truly random ones. The pseudo-random vectors will satisfy (1), (2), and (3) and so they will suffice. This pseudo-randomness and limited independence has deep connections, so we will go into the connections as well.

Error Correcting codes, polynomial interpolation and limited-way independence

Consider the problem of generating a random m -vector $\mathbf{x}$ of ± 1 's so that any four coordinates are mutually independent. We will see that such an m -dimensional vector may be generated from a truly random "seed" of only $O(\log m)$ mutually independent bits. Thus, we need only store the $\log m$ bits and can generate any of the m coordinates when needed. This allows us to store the 4-way independent random m -vector using only $\log m$ bits.

The first fact needed is that for any k , there is a finite field F with exactly 2^k elements, each of which can be represented with k bits and arithmetic operations in the field can be carried out in $O(k^2)$ time. Here, k will be the ceiling of $\log_2 m$. We also assume another basic fact about polynomial interpolation that says that a polynomial of degree at most three is uniquely determined by its value over any field F at four points. More precisely, for any four distinct points a_1, a_2, a_3 , and a_4 in F and any four possibly not distinct values b_1, b_2, b_3 , and b_4 in F , there is a unique polynomial $f(x) = f_0 + f_1x + f_2x^2 + f_3x^3$ of degree at most three, so that with computations done over F , $f(a_1) = b_1, f(a_2) = b_2, f(a_3) = b_3$, and $f(a_4) = b_4$.

Now our definition of the pseudo-random ± 1 vector $\mathbf{x}$ with 4-way independence is simple. Choose four elements f_0, f_1, f_2, f_3 at random from F and form the polynomial $f(s) = f_0 + f_1s + f_2s^2 + f_3s^3$. This polynomial represents $\mathbf{x}$ as follows. For $s = 1, 2, \dots, m$, x_s is the leading bit of the k -bit representation of $f(s)$. Thus, the m -dimensional vector $\mathbf{x}$ requires only $O(k)$ bits where $k = \lceil \log m \rceil$.

Lemma 7.3 *The $\mathbf{x}$ defined above has 4-way independence.*

Proof: Assume that the elements of F are represented in binary using ± 1 instead of the traditional 0 and 1. Let s, t, u , and v be any four coordinates of $\mathbf{x}$ and let $\alpha, \beta, \gamma, \delta$ be any four elements from $\{-1, 1\}$. There are exactly 2^{k-1} elements of F whose leading bit is α and similarly for β, γ , and δ . So, there are exactly $2^{4(k-1)}$ 4-tuples of elements $b_1, b_2, b_3, b_4 \in F$ so that the leading bit of b_1 is α , the leading bit of b_2 is β , the leading bit of b_3 is γ , and the leading bit of b_4 is δ . For each such b_1, b_2, b_3 , and b_4 , there is precisely one polynomial f so that

$$f(s) = b_1, f(t) = b_2, f(u) = b_3, \text{ and } f(v) = b_4.$$

Thus, the probability that the four coordinates x_s, x_t, x_u, x_v of $\mathbf{x}$ are

$$x_s = \alpha, x_t = \beta, x_u = \gamma, \text{ and } x_v = \delta$$

is precisely $\frac{2^{4(k-1)}}{\text{total number of } f} = \frac{2^{4(k-1)}}{2^{4k}} = \frac{1}{16}$ as asserted. ■

The lemma describes how to get one vector $\mathbf{x}$ with 4-way independence. However, we need $r = O(1/\varepsilon^2)$ vectors. Also the vectors must be mutually independent. But this is easy, just choose r polynomials at the outset.

To implement the algorithm with low space, store only the polynomials in memory. This requires $4k = O(\log m)$ bits per polynomial for a total of $O(\log m/\varepsilon^2)$ bits. When a symbol s in the stream is read, compute $x_s^{(1)}, x_s^{(2)}, \dots, x_s^{(r)}$ and update the running sums. Note that x_{s1} is just the leading bit of the first polynomial evaluated at s ; this calculation is in $O(\log m)$ time. Thus, we repeatedly compute the $x_s^{(1)}$ from the “seeds”, namely the

coefficients of the polynomials.

This idea of polynomial interpolation is also used in other contexts. Error-correcting codes is an important example. Say we wish to transmit n bits over a channel which may introduce noise. One can introduce redundancy into the transmission so that some channel errors can be corrected. A simple way to do this is to view the n bits to be transmitted as coefficients of a polynomial $f(x)$ of degree $n - 1$. Now transmit f evaluated at points $1, 2, 3, \dots, n + m$. At the receiving end, any n correct values will suffice to reconstruct the polynomial and the true message. So up to m errors can be tolerated. But even if the number of errors is at most m , it is not a simple matter to know which values are corrupted. We do not elaborate on this here.

7.2 Matrix Algorithms Using Sampling

How does one deal with a large matrix? An obvious suggestion is to take a sample of the matrix. Uniform sampling does not work in general. For example, if a small fraction of the entries are the big/significant ones in the matrix, uniform sampling may miss them all. So the sampling probabilities need to take into account the size or magnitude of the entries. It turns out that sampling the rows and columns of a matrix with probabilities dependent on their length is a good idea in many contexts. We present two examples here, matrix multiplication and the sketch of a matrix. The sketch of a matrix A is a rank r matrix that closely approximates A in the Frobenius norm.

7.2.1 Matrix Multiplication Using Sampling

Suppose A is an $m \times n$ matrix and B is an $n \times p$ matrix and the product AB is desired. We show how to use sampling to get an approximate product faster than traditional matrix multiplication. Let $A(:, k)$ denote the k^{th} column of A . $A(:, k)$ is a $m \times 1$ matrix. Let $B(k, :)$ be the k^{th} row of B . $B(k, :)$ is a $1 \times n$ matrix. It is easy to see that

$$AB = \sum_{k=1}^n A(:, k)B(k, :).$$

Note that for each value of k , $A(:, k)B(k, :)$ is an $m \times p$ matrix each element of which is a single product of elements of A and B . An obvious use of sampling suggests itself. Compute the sum of $A(:, k)B(k, :)$ for some sampled k 's and suitably scale the sum for the estimate of AB .

It turns out that nonuniform sampling probabilities are useful. Define a random variable z that takes on values in $\{1, 2, \dots, n\}$. Let p_k denote the probability that z assumes the value k . The p_k are nonnegative and sum to one. Define an associated random matrix

variable that has value

with probability

$$X = \begin{cases} \frac{1}{p_1} A(:, 1) B(1, :) & p_1 \\ \frac{1}{p_2} A(:, 2) B(2, :) & p_2 \\ \vdots & \\ \frac{1}{p_n} A(:, n) B(n, :) & p_n \end{cases} \quad (7.1)$$

with probability p_k . The matrix X approximates the product AB using only one column of A and one row of B . Let $E(X)$ denote the entry-wise expectation.

$$E(X) = \sum_{k=1}^n \text{Prob}(z = k) \frac{1}{p_k} A(:, k) B(k, :) = \sum_{k=1}^n A(:, k) B(k, :) = AB.$$

This explains the scaling by $\frac{1}{p_k}$ in X .

Define the variance of X as the sum of the variances of all its entries.

$$\begin{aligned} \text{Var}(X) &= \sum_{i=1}^m \sum_{j=1}^p \text{Var}(x_{ij}) = \sum_{ij} E(x_{ij}^2) - \sum_{ij} E^2(x_{ij}) \\ &= \sum_{ij} \sum_k p_k \frac{1}{p_k^2} a_{ik}^2 b_{kj}^2 - \sum_{ij} (AB)_{ij}^2. \end{aligned}$$

Exchanging the order of summations

$$\begin{aligned} \text{Var}(X) &= \sum_k \frac{1}{p_k} \sum_i a_{ik}^2 \sum_j b_{kj}^2 - \sum_{ij} (AB)_{ij}^2 \\ &= \sum_k \frac{1}{p_k} |A(:, k)|^2 |B(k, :)|^2 - \sum_{ij} (AB)_{ij}^2. \end{aligned}$$

What is the best choice of p_k ? It is the one which minimizes the variance. The term $\sum_{ij} (AB)_{ij}^2 = \|AB\|_F^2$. So we should choose p_k to minimize $\sum_k \frac{1}{p_k} |A(:, k)|^2 |B(k, :)|^2$. It can be seen by calculus that the minimizing p_k are proportional to $|A(:, k)| |B(k, :)|$. In the important special case when $B = A^T$, one should pick columns of A with probabilities proportional to the squared length of the columns.

In the general case when B is not A^T , length squared sampling simplifies bounds. If p_k is proportional to $|A(:, k)|^2$, i.e, $p_k = \frac{|A(:, k)|^2}{\|A\|_F^2}$, we can bound $\text{Var}(X)$ by

$$\text{Var}(X) \leq \|A\|_F^2 \sum_k |B(k, :)|^2 = \|A\|_F^2 \|B\|_F^2.$$

To reduce the variance, do s independent trials. Each trial i , $i = 1, 2, \dots, s$ yields a matrix X_i as in (7.1). Take $\frac{1}{s} \sum_{i=1}^s X_i$ as our estimate of AB . Since the variance of a

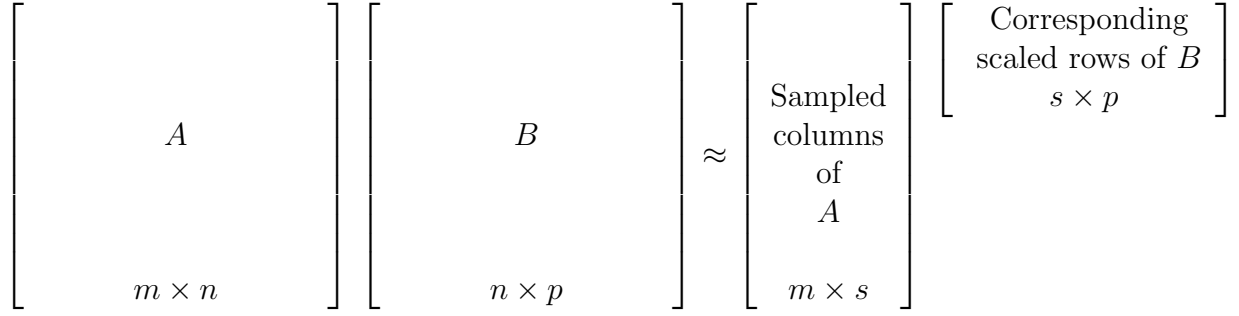


Figure 7.2: Approximate Matrix Multiplication using sampling

sum of independent random variables is the sum of variances, the variance of $\frac{1}{s} \sum_{i=1}^s X_i$ is $\frac{1}{s} \text{Var}(X)$ and is at most $\frac{1}{s} \|A\|_F^2 \|B\|_F^2$.

To implement this, suppose $k_1, k_2, \dots, k_s$ are the k 's chosen in each trial. It is easy to see that

$$\frac{1}{s} \sum_{i=1}^s X_i = \frac{1}{s} \left(\frac{A(:, k_1) B(k_1, :)}{p_{k_1}} + \frac{A(:, k_2) B(k_2, :)}{p_{k_2}} + \dots + \frac{A(:, k_s) B(k_s, :)}{p_{k_s}} \right) = C \tilde{B},$$

where, C is the $m \times s$ matrix of the chosen columns of A and $\tilde{B}$ is an $s \times p$ matrix with the corresponding rows of B scaled, namely, $\tilde{B}$ has rows $B(k_1, :)/(sp_{k_1})$, $B(k_2, :)/(sp_{k_2})$, $\dots$, $B(k_s, :)/(sp_{k_s})$. See Figure (7.2).

We summarize our discussion in a lemma.

Lemma 7.4 *Suppose A is an $m \times n$ matrix and B is an $n \times p$ matrix. The product AB can be estimated by $C\tilde{B}$, where, C is an $m \times s$ matrix consisting of s columns of A picked in independent trials, each according to length-squared distribution and $\tilde{B}$ is the $s \times p$ matrix consisting of the corresponding rows of B scaled as above. The error is bounded by:*

$$E \left(\|AB - C\tilde{B}\|_F^2 \right) \leq \frac{\|A\|_F^2 \|B\|_F^2}{s}.$$

7.2.2 Sketch of a Large Matrix

The main result of this section will be that for any matrix, a sample of columns and rows, each picked in independent trials according to length squared distribution is a good sketch of the matrix.

Let A be an $m \times n$ matrix. Pick s columns of A in independent trials, in each picking a column according to length squared distribution on the columns. Let C be the $m \times s$ matrix containing the picked columns. Similarly, pick r rows of A in r independent trials,

each according to length squared distribution on the rows of A . Let R be the $r \times n$ matrix of the picked rows. From C and R , we can compute a matrix U so that $A \approx CUR$. The schematic diagram is given in Figure 7.3.

The proof makes crucial use of the fact that the sampling of rows and columns is with probability proportional to the squared length. One may recall that the top k singular vectors of the SVD of A , give a similar picture; but the SVD takes more time to compute, requires all of A to be stored in RAM, and does not have the property that the rows and columns are directly from A . The last property, that the approximation involves actual rows and columns of the matrix rather than linear combinations, is called an *interpolative approximation* and is useful in many contexts. However, the SVD does yield the best 2-norm approximation. Error bounds for the approximation CUR are weaker.

We briefly touch upon two motivations for such a sketch. Suppose A is the document-term matrix of a large collection of documents. We are to “read” the collection at the outset and store a sketch so that later, when a query represented by a vector with one entry per term arrives, we can find its similarity to each document in the collection. Similarity is defined by the dot product. In Figure 7.3 it is clear that the matrix-vector product of a query with the right hand side can be done in time $O(ns + sr + rm)$ which would be linear in n and m if s and r are $O(1)$. The error bound for this process, requires that the difference between A and the sketch of A has small 2-norm. Recall that the 2-norm $\|A\|_2$ of a matrix A is $\max_{\|\mathbf{x}\|=1} |A\mathbf{x}|$. The fact that the sketch is an interpolative approximation means that our approximation essentially consists of a subset of documents and a subset of terms, which may be thought of as a representative set of documents and terms.

A second motivation comes from recommendation systems. Here A would be a customer-product matrix whose $(i, j)^{th}$ entry is the preference of customer i for product j . The objective is to collect a few sample entries of A and based on them, get an approximation to A so that we can make future recommendations. A few sampled rows of A (all preferences of a few customers) and a few sampled columns (all customers’ preferences for a few products) give a good approximation to A provided that the samples are drawn according to the length-squared distribution.

It remains to describe how to find U from C and R . Through the rest of this section, we make the assumption that RR^T is invertible. This case will convey the essential ideas. Also, note that since r in general will be much smaller than n and m , unless the matrix A is degenerate, it is likely that the r rows in the sample R will be linearly independent giving us invertibility of RR^T .

Before stating precisely what U is, we give some intuition. Write A as AI , where, I is the $n \times n$ identity matrix. Pretend for the moment that we approximate the product AI by sampling s columns of A according to length-squared. Then, as in the last section, write $AI \approx CW$ where, W consists of a scaled version of the s rows of I cor-

sampling and similarly R is a matrix of r rows of A picked according to length squared sampling. Then, we can find from C, R an $s \times r$ matrix U so that

$$E(\|A - CUR\|_2^2) \leq \|A\|_F^2 \left(\frac{2}{r} + \frac{2r}{s} \right).$$

Choosing $s = r^2$, the bound becomes $O(1/r)\|A\|_F^2$ and if want the bound to be at most $\varepsilon\|A\|_F^2$ for some small $\varepsilon > 0$, it suffices to choose $r \in \Omega(1/\varepsilon)$.

We now briefly look at the time needed to compute U . The only involved step in computing U is to find $(RR^T)^{-1}$. But note that RR^T is an $s \times s$ matrix and since s is to much smaller than n and m , this is fast. Now we prove all the claims used in the discussion above.

Lemma 7.8 *If RR^T is invertible, then $R^T(RR^T)^{-1}R$ acts as the identity matrix on the row space of R . I.e., for every vector $\mathbf{x}$ of the form $\mathbf{x} = R^T\mathbf{y}$ (this defines the row space of R), we have $R^T(RR^T)^{-1}R\mathbf{x} = \mathbf{x}$.*

Proof: For $\mathbf{x} = R^T\mathbf{y}$, since RR^T is invertible

$$R^T(RR^T)^{-1}R\mathbf{x} = R^T(RR^T)^{-1}RR^T\mathbf{y} = R^T\mathbf{y} = \mathbf{x}$$

■

Now we prove Proposition 7.5. First suppose $\mathbf{x} \in V$. Then we can write $\mathbf{x} = R^T\mathbf{y}$ and so $P\mathbf{x} = R^T(RR^T)^{-1}RR^T\mathbf{y} = R^T\mathbf{y} = \mathbf{x}$, so for $\mathbf{x} \in V$, we have $(A - AP)\mathbf{x} = \mathbf{0}$. So, it suffices to consider $\mathbf{x} \in V^\perp$. For such $\mathbf{x}$, $(A - AP)\mathbf{x} = A\mathbf{x}$ and

$$|(A - AP)\mathbf{x}|^2 = |A\mathbf{x}|^2 = \mathbf{x}^T A^T A \mathbf{x} = \mathbf{x}^T (A^T A - R^T R) \mathbf{x} \leq \|A^T A - R^T R\|_2 |\mathbf{x}|^2,$$

so we get $\|A - AP\|_2^2 \leq \|A^T A - R^T R\|_2$, so it suffices to prove that $\|A^T A - R^T R\|_2 \leq \|A\|_F^2/r$ which follows directly from Lemma 7.4 since we can think of $R^T R$ as a way of estimating $A^T A$ by picking (according to length-squared distribution) columns of A^T , i.e., rows of A . This proves Proposition 7.5.

Proposition 7.6 is easy to see: Since by Lemma 7.8, P is the identity on the space V spanned by the rows of R , we have that $\|P\|_F^2$ is the sum of its singular values squared which is at most r as claimed.

7.3 Sketches of Documents

Suppose one wished to store all the web pages from the WWW. Since there are billions of web pages, one might store just a sketch of each page where a sketch is a few hundred bits that capture sufficient information to do whatever task one had in mind. A web page

or a document is a sequence. We first show how to sample a set and then how to convert the problem of sampling a sequence into the problem of sampling a set.

Consider subsets of size 1000 of the integers from 1 to 10^6 . Suppose one wished to compute the resemblance of two subsets A and B by the formula

$$\text{resemblance}(A, B) = \frac{|A \cap B|}{|A \cup B|}$$

Suppose that instead of using the sets A and B , one sampled the sets and compared random subsets of size ten. How accurate would the estimate be? One way to sample would be to select ten elements uniformly at random from A and B . However, this method is unlikely to produce overlapping samples. Another way would be to select the ten smallest elements from each of A and B . If the sets A and B overlapped significantly one might expect the sets of ten smallest elements from each of A and B to also overlap. One difficulty that might arise is that the small integers might be used for some special purpose and appear in essentially all sets and thus distort the results. To overcome this potential problem, rename all elements using a random permutation.

Suppose two subsets of size 1000 overlapped by 900 elements. What would the overlap be of the 10 smallest elements from each subset assuming that the elements have been renamed using a random permutation? One would expect the nine smallest elements from the 900 common elements to be in each of the two subsets for an overlap of 90%. The expected resemblance for the size ten sample would be $9/11=0.81$, which is the $\text{resemblance}(A, B)$.

Another method would be to select the elements equal to zero mod m for some integer m . If one samples mod m , the size of the sample becomes a function of n and thus sampling mod m allows us to also handle containment.

In another version of the problem, one has a sequence rather than a set. Here one converts the sequence into a set by replacing the sequence by the set of all short subsequences of some length k . Corresponding to each sequence is a set of length k subsequences. If k is sufficiently large, then two sequences are highly unlikely to give rise to the same set of subsequences. Thus, we have converted the problem of sampling a sequence to that of sampling a set. Instead of storing all the subsequences, one stores only a small subset of the set of length k subsequences.

Suppose you wish to be able to determine if two web pages are minor modifications of one another or to determine if one is a fragment of the other. Define the set of subsequences of k consecutive words from the sequence of words on the page. Let $S(D)$ be the set of all subsequences of length k occurring in document D . Define resemblance of A and B by

$$\text{resemblance}(A, B) = \frac{|S(A) \cap S(B)|}{|S(A) \cup S(B)|}$$

And define containment as

$$\text{containment}(A, B) = \frac{|S(A) \cap S(B)|}{|S(A)|}$$

Let W be a set of subsequences. Define $\min(W)$ to be the s smallest elements in W and define $\text{mod}(W)$ as the set of elements of w that are zero mod m .

Let π be a random permutation of all length k subsequences. Define $F(A)$ to be the s smallest elements of A and $V(A)$ to be the set mod m in the ordering defined by the permutation.

Then

$$\frac{F(A) \cap F(B)}{F(A) \cup F(B)}$$

and

$$\frac{|V(A) \cap V(B)|}{|V(A) \cup V(B)|}$$

are unbiased estimates of the resemblance of A and B . The value

$$\frac{|V(A) \cap V(B)|}{|V(A)|}$$

is an unbiased estimate of the containment of A in B .

7.4 Exercises

Exercise 7.1 Given a stream of n positive real numbers $a_1, a_2, \dots, a_n$, upon seeing $a_1, a_2, \dots, a_i$ keep track of the sum $s = a_1 + a_2 + \dots + a_i$ and a sample a_j , $j \leq i$ drawn with probability proportional to its value. On reading a_{i+1} , with probability $\frac{a_{i+1}}{s+a_{i+1}}$ replace the current sample with a_{i+1} and update s . Prove that the algorithm selects an a_j from the stream with the probability of picking a_j being proportional to its value.

Exercise 7.2 Given a stream of symbols $a_1, a_2, \dots, a_n$, give an algorithm that will select one symbol uniformly at random from the stream. How much memory does your algorithm require?

Exercise 7.3 Give an algorithm to select an a_i from a stream of symbols $a_1, a_2, \dots, a_n$ with probability proportional to a_i^2 .

Exercise 7.4 How would one pick a random word from a very large book where the probability of picking a word is proportional to the number of occurrences of the word in the book?

Exercise 7.5 For the streaming model give an algorithm to draw s independent samples each with the probability proportional to its value. Justify that your algorithm works correctly.

Exercise 7.6 Show that for a 2-universal hash family $\text{Prob}(h(x) = z) = \frac{1}{M+1}$ for all $x \in \{1, 2, \dots, m\}$ and $z \in \{0, 1, 2, \dots, M\}$.

Exercise 7.7 Let p be a prime. A set of hash functions

$$H = \{h \mid \{0, 1, \dots, p-1\} \rightarrow \{0, 1, \dots, p-1\}\}$$

is 3-universal if for all u, v, w, x, y , and z in $\{0, 1, \dots, p-1\}$

$$\text{Prob}(h(x) = u, h(y) = v, h(z) = w) = \frac{1}{p^3}.$$

(a) Is the set $\{h_{ab}(x) = ax + b \bmod p \mid 0 \leq a, b < p\}$ of hash functions 3-universal?

(b) Give a 3-universal set of hash functions.

Exercise 7.8 Give an example of a set of hash functions that is not 2-universal.

Exercise 7.9

(a) What is the variance of the method in Section 7.1.2 of counting the number of occurrences of a 1 with $\log \log n$ memory?

(b) Can the algorithm be iterated to use only $\log \log \log n$ memory? What happens to the variance?

Exercise 7.10 Prove that for independent x and y , $\text{Var}(x + y) = \text{Var}(x) + \text{Var}(y)$. *Hint: Shifting a probability distribution so that it has expected value does not change the Var of the distribution.*

Exercise 7.11 Consider a coin that comes down heads with probability p . Prove that the expected number of flips before a head occurs is $1/p$.

Exercise 7.12 Randomly generate a string $x_1x_2\cdots x_n$ of 10^6 0's and 1's with probability $1/2$ of x_i being a 1. Count the number of ones in the string and also estimate the number of ones by the approximate counting algorithm. Repeat the process for $p=1/4$, $1/8$, and $1/16$. How close is the approximation?

Exercise 7.13 Construct an example in which the majority algorithm gives a false positive, i.e., stores a non majority element at the end.

Exercise 7.14 Construct examples where the frequent algorithm in fact does as badly as in the theorem, i.e., it “under counts” some item by $n/(k+1)$.

Exercise 7.15 Recall basic statistics on how an average of independent trials cuts down variance and complete the argument for relative error ε estimate of $\sum_{s=1}^m f_s^2$.

Exercise 7.16 What are the groups $GF(2)$, $GF(4)$, and $GF(8)$? Specify the elements and the operations of addition and multiplication.

Exercise 7.17 Let F be a field. Prove that for any four distinct points a_1, a_2, a_3 , and a_4 in F and any four (possibly not distinct) values b_1, b_2, b_3 , and b_4 in F , there is a unique polynomial $f(x) = f_0 + f_1x + f_2x^2 + f_3x^3$ of degree at most three so that $f(a_1) = b_1$, $f(a_2) = b_2$, $f(a_3) = b_3$, $f(a_4) = b_4$ with all computations done over F .

Exercise 7.18 Suppose we want to pick a row of a matrix at random where the probability of picking row i is proportional to the sum of squares of the entries of that row. How would we do this in the streaming model? Do not assume that the elements of the matrix are given in row order.

- (a) Do the problem when the matrix is given in column order.
- (b) Do the problem when the matrix is represented in sparse notation: it is just presented as a list of triples (i, j, a_{ij}) , in arbitrary order.

Exercise 7.19 Suppose A and B are two matrices. Show that $AB = \sum_{k=1}^n A(:, k)B(k, :)$.

Exercise 7.20 Generate two 100 by 100 matrices A and B with integer values between 1 and 100. Compute the product AB both directly and by sampling. Plot the difference in L_2 norm between the results as a function of the number of samples. In generating the matrices make sure that they are skewed. One method would be the following. First generate two 100 dimensional vectors a and b with integer values between 1 and 100. Next generate the i^{th} row of A with integer values between 1 and a_i and the i^{th} column of B with integer values between 1 and b_i .

Exercise 7.21 Show that $ADD^T B$ is exactly

$$\frac{1}{s} \left(\frac{A(:, k_1) B(k_1, :)}{p_{k_1}} + \frac{A(:, k_2) B(k_2, :)}{p_{k_2}} + \dots + \frac{A(:, k_s) B(k_s, :)}{p_{k_s}} \right)$$

Exercise 7.22 Suppose $a_1, a_2, \dots, a_n$ are nonnegative reals. Show that the minimum of $s = \sum_{k=1}^n \frac{a_k^2}{x_k}$ subject to the constraints $x_k \geq 0$ and $\sum_{k=1}^n x_k = 1$ is attained when the x_k are proportional to a_k .

Solution: $\frac{\delta s}{\delta p_i} = -\frac{a_i^2}{p_i^2}$. If the i^{th} derivative decreases s more increasing p_i than the j^{th} derivative increases s by decreasing p_j , then s can be reduced. Thus, all derivatives are equal at the minimum. Thus, $p_i = \frac{a_i}{\sum_{j=1}^n a_j}$ for all i when s is minimum. ■

Exercise 7.23 Consider random sequences of length n composed of the integers 0 through 9. Represent a sequence by its set of length k -subsequences. What is the resemblance of the sets of length k -subsequences from two random sequences of length n for various values of k as n goes to infinity?

Exercise 7.24 What if the sequences in the Exercise 7.23 were not random? Suppose the sequences were strings of letters and that there was some nonzero probability of a given letter of the alphabet following another. Would the result get better or worse?

Exercise 7.25 Consider a random sequence of length 10,000 over an alphabet of size 100.

- (a) For $k = 3$ what is probability that two possible successor subsequences for a given subsequence are in the set of subsequences of the sequence?
- (b) For $k = 5$ what is the probability?

Exercise 7.26 How would you go about detecting plagiarism in term papers?

Exercise 7.27 Suppose you had one billion web pages and you wished to remove duplicates. How would you do this?

Exercise 7.28 Construct two sequences of 0's and 1's having the same set of subsequences of width w .

Exercise 7.29 Consider the following lyrics:

When you walk through the storm hold your head up high and don't be afraid of the dark. At the end of the storm there's a golden sky and the sweet silver song of the lark.

Walk on, through the wind, walk on through the rain though your dreams be tossed and blown. Walk on, walk on, with hope in your heart and you'll never walk alone, you'll never walk alone.

How large must k be to uniquely recover the lyric from the set of all subsequences of symbols of length k ? Treat the blank as a symbol.

Exercise 7.30 ***Blast:** Given a long sequence a , say 10^9 and a shorter sequence b , say 10^5 , how do we find a position in a which is the start of a subsequence b' that is close to b ? This problem can be solved by dynamic programming but not in reasonable time. Find a time efficient algorithm to solve this problem.*

*Hint: (**Shingling approach**) One possible approach would be to fix a small length, say seven, and consider the shingles of a and b of length seven. If a close approximation to b is a substring of a , then a number of shingles of b must be shingles of a . This should allow us to find the approximate location in a of the approximation of b . Some final algorithm should then be able to find the best match.*

8 Clustering

8.1 Some Clustering Examples

Clustering refers to the process of partitioning a set of objects into subsets consisting of similar objects. Clustering comes up in many contexts. For example, one might want to cluster journal articles into clusters of articles on related topics. In doing this, one first represents a document by a vector. This can be done using the vector space model introduced in Chapter 2. Each document is represented as a vector with one component for each term giving the frequency of the term in the document. Alternatively, a document may be represented by a vector whose components correspond to documents in the collection and the j^{th} component of the i^{th} vector is a 0 or 1 depending on whether the i^{th} document referenced the j^{th} document. Once one has represented the documents as vectors, the problem becomes one of clustering vectors.

Another context where clustering is important is the study of the evolution and growth of communities in social networks. Here one constructs a graph where nodes represent individuals and there is an edge from one node to another if the person corresponding to the first node sent an email or instant message to the person corresponding to the second node. A community is defined as a set of nodes where the frequency of messages within the set is higher than what one would expect if the set of nodes in the community were a random set. Clustering partitions the set of nodes of the graph into sets of nodes where the sets consist of nodes that send more messages to one another than one would expect by chance. Note that clustering generally asks for a strict partition into subsets, although in reality a node may belong to several communities.

In these clustering problems, one defines either a similarity measure between pairs of objects or a distance measure, a notion of dissimilarity. One measure of similarity between two vectors $\mathbf{a}$ and $\mathbf{b}$ is the cosine of the angle between them:

$$\cos(\mathbf{a}, \mathbf{b}) = \frac{\mathbf{a}^T \mathbf{b}}{|\mathbf{a}| |\mathbf{b}|}.$$

To get a distance measure, subtract the cosine similarity from one.

$$\text{dist}(\mathbf{a}, \mathbf{b}) = 1 - \cos(\mathbf{a}, \mathbf{b})$$

Another distance measure is the Euclidean distance. There is an obvious relationship between cosine similarity and Euclidean distance. If $\mathbf{a}$ and $\mathbf{b}$ are unit vectors, then

$$|\mathbf{a} - \mathbf{b}|^2 = (\mathbf{a} - \mathbf{b})^T (\mathbf{a} - \mathbf{b}) = |\mathbf{a}|^2 + |\mathbf{b}|^2 - 2\mathbf{a}^T \mathbf{b} = 2(1 - \cos(\mathbf{a}, \mathbf{b})).$$

In determining the distance function to use, it is useful to know something about the origin of the data. In clustering the nodes of a graph, we may represent each node as a vector, namely, as the row of the adjacency matrix corresponding to the node. One

notion of dissimilarity here is the square of the Euclidean distance. For 0-1 vectors, this measure is just the number of “uncommon” 1’s, whereas, the dot product is the number of common 1’s.

In many situations one has a stochastic model of how the data was generated. An example is customer behavior. Suppose there are d products and n customers. A reasonable assumption is that each customer generates from a probability distribution, the basket of goods he or she buys. A basket specifies the amount of each good bought. One hypothesis is that there are only k types of customers, $k \ll n$. Each customer type is characterized by a probability density used by all customers of that type to generate their baskets of goods. The densities may all be Gaussians with different centers and covariance matrices. We are not given the probability densities, only the basket bought by each customer, which is observable. Our task is to cluster the customers into the k types. We may identify the customer with his or her basket which is a vector. One way to formulate the problem mathematically is by a clustering criterion that is then optimized. Some potential criteria are to partition the customers into k clusters so as to minimize

1. the sum of distances between all pairs of customers in the same cluster,
2. the sum of distances of all customers to their “cluster center” (any point in space may be designated as the cluster center), or
3. the sum of squared distances to the cluster center.

The last criterion is called the *k-means* criterion and is widely used. The variant (2) above called the *k-median* criterion minimizes the sum of distances (not squared) to the cluster center. Another possibility, called the *k-center* criterion, is to minimize the maximum distance of any point to its cluster center.

The chosen criterion can affect the results. To illustrate, suppose that the data was generated according to an equal weight mixture of k spherical Gaussian densities centered at $\boldsymbol{\mu}_1, \boldsymbol{\mu}_2, \dots, \boldsymbol{\mu}_k$, each with variance one in every direction. Then the density of the mixture is

$$F(\mathbf{x}) = \text{Prob}(\mathbf{x}) = \frac{1}{k} \frac{1}{(2\pi)^{d/2}} \sum_{i=1}^k e^{-|\mathbf{x}-\boldsymbol{\mu}_i|^2}.$$

Denote by $\boldsymbol{\mu}(\mathbf{x})$ the center nearest to $\mathbf{x}$. Since the exponential function falls off fast, we can approximate $\sum_{i=1}^k e^{-|\mathbf{x}-\boldsymbol{\mu}_i|^2}$ by $e^{-|\mathbf{x}-\boldsymbol{\mu}(\mathbf{x})|^2}$. Thus

$$F(\mathbf{x}) \approx \frac{1}{k} \frac{1}{(2\pi)^{d/2}} e^{-|\mathbf{x}-\boldsymbol{\mu}(\mathbf{x})|^2}.$$

The likelihood of drawing the sample of points $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(n)}$ from the mixture, if the centers were $\boldsymbol{\mu}_1, \boldsymbol{\mu}_2, \dots, \boldsymbol{\mu}_k$, is approximately

$$\frac{1}{k^n} \frac{1}{(2\pi)^{nd/2}} \prod_{i=1}^n e^{-|\mathbf{x}^{(i)}-\boldsymbol{\mu}(\mathbf{x}^{(i)})|^2} = c e^{-\sum_{i=1}^n |\mathbf{x}^{(i)}-\boldsymbol{\mu}(\mathbf{x}^{(i)})|^2}.$$

Minimizing the sum of squared distances to cluster centers finds the maximum likelihood $\mu_1, \mu_2, \dots, \mu_k$. This suggests using the sum of distance squared to the cluster centers.

On the other hand, if the generating process had an exponential probability distribution, with the probability law

$$\text{Prob}[(x_1, x_2, \dots, x_d)] = \frac{1}{2^d} \prod_{i=1}^d e^{-|x_i - \mu_i|} = \frac{1}{2^d} e^{-\sum_{i=1}^d |x_i - \mu_i|} = \frac{1}{2^d} e^{-\|\mathbf{x} - \boldsymbol{\mu}\|_1},$$

one would use the L_1 norm, not the L_2 or the square of the L_1 , since the probability density decreases as the L_1 distance from the center. The intuition here is that the distance used to cluster data should be related to the actual distribution of the data.

The choice of whether to use a distance measure and cluster together points that are close or use a similarity measure and cluster together points with high similarity, and what particular distance or similarity measure to use, can be crucial to the application. However, there is not much theory on these choices; they are determined by empirical domain-specific knowledge. One general observation is worth making. Using distance squared instead of distance, favors outliers since the square function magnifies large values, which means a small number of outliers may make a clustering look bad. On the other hand, distance squared has some mathematical advantages; see for example Corollary 8.2 that asserts that with the distance squared criterion, the centroid is the correct cluster center. The widely used k -means criterion is based on sum of squared distances.

There are in general two variations of the clustering problem for each of the criteria. We could require that each cluster center be a data point or allow a cluster center to be any point in space. If we require each cluster center to be a data point, the optimal clustering of n data points into k clusters can be solved in time $\binom{n}{k}$ times a polynomial in the length of the data. First, exhaustively enumerate all sets of k data points as the possible sets of k cluster centers, then associate each point to its nearest center and select the best clustering. No such naive enumeration procedure is available when cluster centers can be any point in space. But, for the k -means problem, Corollary 8.2 shows that once we have identified the data points that belong to a cluster, the best choice of cluster center is the centroid. Note that the centroid might not be a data point.

In the formulations discussed so far, we have one number (e.g. sum of distances squared to the cluster center) as the measure of goodness of a clustering and we try to optimize that number to find the best clustering according to the measure. This approach does not always yield desired results, since it can be hard to find the optimum exactly. Although most clustering problems are NP-hard, often there are polynomial time algorithms to find an approximately optimal solution. But such a solution may be far from the optimal or desired clustering. We will see in Section 8.4 how to formalize some realistic conditions under which an approximate optimal solution gives us a desired clustering as well. But

first we see some simple algorithms for getting a good clustering according to some natural measures.

8.2 A k -means Clustering Algorithm

There are many algorithms for clustering high dimensional data. We start with a widely used algorithm that uses the k -means criterion. In the k -means criterion, a set $A = \{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n\}$ of n points in d -dimensions is partitioned into k -clusters, $S_1, S_2, \dots, S_k$, so as to minimize the sum of squared distances of each point to its cluster center. That is, A is partitioned into clusters, $S_1, S_2, \dots, S_k$, and a center is assigned to each cluster so as to minimize

$$d(S_1, S_2, \dots, S_k) = \sum_{j=1}^k \sum_{\mathbf{a}_i \in S_j} (\mathbf{c}_j - \mathbf{a}_i)^2$$

where $\mathbf{c}_j$ is the center of cluster j .

Suppose we have already determined the clustering or the partitioning into $S_1, S_2, \dots, S_k$. What are the best centers for the clusters? The following lemma shows that the answer is the centroids, the coordinate means, of the clusters.

Lemma 8.1 *Let $\{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n\}$ be a set of points. The sum of the squared distances of the $\mathbf{a}_i$ to any point $\mathbf{x}$ equals the sum of the squared distances to the centroid plus the number of points times the squared distance from the point $\mathbf{x}$ to the centroid. That is,*

$$\sum_i |\mathbf{a}_i - \mathbf{x}|^2 = \sum_i |\mathbf{a}_i - \mathbf{c}|^2 + n |\mathbf{c} - \mathbf{x}|^2$$

where $\mathbf{c} = \frac{1}{n} \sum_{i=1}^n \mathbf{a}_i$ is the centroid of the set of points.

Proof:

$$\begin{aligned} \sum_i |\mathbf{a}_i - \mathbf{x}|^2 &= \sum_i |\mathbf{a}_i - \mathbf{c} + \mathbf{c} - \mathbf{x}|^2 \\ &= \sum_i |\mathbf{a}_i - \mathbf{c}|^2 + 2(\mathbf{c} - \mathbf{x}) \cdot \sum_i (\mathbf{a}_i - \mathbf{c}) + n |\mathbf{c} - \mathbf{x}|^2 \end{aligned}$$

Since $\mathbf{c}$ is the centroid, $\sum_i (\mathbf{a}_i - \mathbf{c}) = 0$. Thus, $\sum_i |\mathbf{a}_i - \mathbf{x}|^2 = \sum_i |\mathbf{a}_i - \mathbf{c}|^2 + n |\mathbf{c} - \mathbf{x}|^2$ ■

A corollary of Lemma 8.1 is that the centroid minimizes the sum of squared distances since the second term, $n \|\mathbf{c} - \mathbf{x}\|^2$, is always non-negative.

Corollary 8.2 *Let $\{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n\}$ be a set of points. The sum of squared distances of the $\mathbf{a}_i$ to a point $\mathbf{x}$ is minimized when $\mathbf{x}$ is the centroid, namely $\mathbf{x} = \frac{1}{n} \sum_i \mathbf{a}_i$.*

Another expression for the sum of squared distances of a set of n points to their centroid is the sum of all pairwise distances squared divided by n . First, a simple observation. For a set of points $\{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n\}$, $\sum_{i=1}^n \sum_{j=i+1}^n |\mathbf{a}_i - \mathbf{a}_j|^2$ counts the quantity $|\mathbf{a}_i - \mathbf{a}_j|^2$ once for each ordered pair (i, j) , $j > i$. However, $\sum_{i,j} |\mathbf{a}_i - \mathbf{a}_j|^2$ counts each $|\mathbf{a}_i - \mathbf{a}_j|^2$ twice, so the later sum is twice the first sum.

Lemma 8.3 *Let $\{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n\}$ be a set of points. The sum of the squared distances between all pairs of points equals the number of points times the sum of the squared distances of the points to the centroid of the points. That is, $\sum_i \sum_{j>i} |\mathbf{a}_i - \mathbf{a}_j|^2 = n \sum_i |\mathbf{a}_i - \mathbf{c}|^2$ where $\mathbf{c}$ is the centroid of the set of points.*

Proof: Lemma 8.1 states that for every $\mathbf{x}$,

$$\sum_i |\mathbf{a}_i - \mathbf{x}|^2 = \sum_i |\mathbf{a}_i - \mathbf{c}|^2 + n |\mathbf{c} - \mathbf{x}|^2.$$

Letting $\mathbf{x}$ range over all $\mathbf{a}_j$ and summing the n equations yields

$$\begin{aligned} \sum_{i,j} |\mathbf{a}_i - \mathbf{a}_j|^2 &= n \sum_i |\mathbf{a}_i - \mathbf{c}|^2 + n \sum_j |\mathbf{c} - \mathbf{a}_j|^2 \\ &= 2n \sum_i |\mathbf{a}_i - \mathbf{c}|^2. \end{aligned}$$

Observing that

$$\sum_{i,j} |\mathbf{a}_i - \mathbf{a}_j|^2 = 2 \sum_i \sum_{j>i} |\mathbf{a}_i - \mathbf{a}_j|^2$$

yields the result that

$$\sum_i \sum_{j>i} |\mathbf{a}_i - \mathbf{a}_j|^2 = n \sum_i |\mathbf{a}_i - \mathbf{c}|^2.$$

■

The k -means clustering algorithm

A natural algorithm for k -means clustering is given below. There are three unspecified aspects of the algorithm. One is k , the number of clusters, a second is the actual set of starting centers and the third is the stopping condition.

The k -means algorithm

Start with k centers.

Cluster each point with the center nearest to it.

Find the centroid of each cluster and replace the set of old centers with the centroids.

Repeat the above two steps until the centers converge (according to some criterion).

The k -means algorithm always converges but often to a local minimum. To show convergence, we argue that the sum of the squares of the distances of each point to its cluster center, always improves. Each iteration consists of two steps. First, consider the step that finds the centroid of each cluster and replaces the old centers with the new centers. By Corollary 8.2, this step improves the sum of internal cluster distances squared. The second step reclusters by assigning each point to its nearest cluster center, which also improves the internal cluster distances.

One way to determine a good value of k is to run the algorithm for each value of k and plot the sum of squared distances to the cluster centers as a function of k . If the value of the sum drops sharply going from some value of k to $k + 1$, then this suggests that $k + 1$ corresponds to the number of clusters in a natural partition of the data.

Another issue that arises is whether the clusters have any real significance. The k -means algorithm will find k clusters even in $G(n, p)$. But note that since the graph $G(n, p)$ should look uniform everywhere, there aren't really k meaningful clusters where a clustering is meaningful if any close to optimal clustering is almost identical to it. In fact, there are many ways of clustering the vertices of this graph all of which will be nearly optimal with respect to the k -means or k -median criteria.

8.3 A Greedy Algorithm for k -Center Criterion Clustering

In this section, instead of using the k -means clustering criterion, we use the k -center criterion. The k -center criterion partitions the points into k clusters so as to minimize the maximum distance of any point to its cluster center. Call the maximum distance of any point to its cluster center the radius of the clustering. There is a k -clustering of radius r if and only if there are k spheres, each of radius r , which together cover all the points. Below, we give a simple algorithm to find k spheres covering a set of points. The following lemma shows that this algorithm only needs to use a radius that is “off by a factor of at most two” from the optimal k -center solution.

The Greedy k -clustering Algorithm

Pick any data point to be the first cluster center. At time t , for $t = 2, 3, \dots, k$, pick any data point that is not within distance r of an existing cluster center; make it the t^{th} cluster center.

Lemma 8.4 *If there is a k -clustering of radius $\frac{r}{2}$, then the above algorithm finds a k -clustering with radius at most r .*

Proof: Suppose for contradiction that the algorithm using radius r fails to find a k -clustering. This means that after the algorithm chooses k centers, there is still at least one data point that is not in any sphere of radius r around a picked center. This is the only possible mode of failure. But then there are $k + 1$ data points, with each pair more

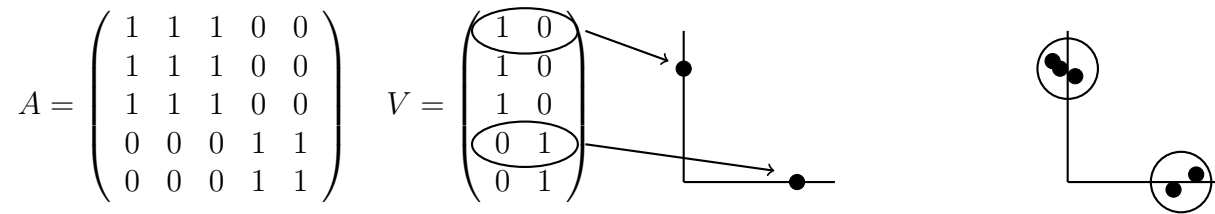


Figure 8.1: Illustration of spectral clustering.

than distance r apart. Clearly, no two such points can belong to the same cluster in any k -clustering of radius $\frac{r}{2}$ contradicting the hypothesis. ■

8.4 Spectral Clustering

In this section we give two contexts where spectral clustering is used. The first is used for finding communities in graphs and the second for clustering a general set of data points. We begin with a simple explanation as to how spectral clustering works when applied to the adjacency matrix of a graph.

Spectral clustering applied to graphs

In spectral clustering of the vertices of a graph, one first creates a new matrix V whose columns correspond to the first k singular vectors of the adjacency matrix. Each row of V is the projection of a row of the adjacency matrix to the space spanned by the k singular vectors. In the example below, the graph has five vertices divided into two cliques, one consisting of the first three vertices and the other the last two vertices. The top two right singular vectors of the adjacency matrix, not normalized to length one, are $(1, 1, 1, 0, 0)^T$ and $(0, 0, 0, 1, 1)^T$. The five rows of the adjacency matrix projected to these vectors form the 5×2 matrix in Figure 8.1. Here, in fact there are two ideal clusters with all edges inside a cluster being present including all self-loops and all edges between clusters being absent. The five rows project to just two points, depending on which cluster the rows are in. If the clusters were not so ideal and instead of the graph consisting of two disconnected cliques, the graph consisted of two dense subsets of vertices where the two sets were connected by only a few edges, then the singular vectors would not be indicator vectors for the clusters but close to indicator vectors. The rows would be mapped to two clusters of points instead of two points. A k -means clustering algorithm would find the clusters.

If the clusters were overlapping, then instead of two clusters of points, there would be

three clusters of points where the third cluster corresponds to the overlapping vertices of the two clusters. Instead of using k -means clustering, we might instead find the minimum 1-norm vector in the space spanned by the two singular vectors. The minimum 1-norm vector will not be an indicator vector, so we would threshold its values to create an indicator vector for a cluster. Instead of finding the minimum 1-norm vector in the space spanned by the singular vectors in V , we might actually look for a small 1-norm vector close to the subspace.

$$\min_{\mathbf{x}} (1 - |\mathbf{x}|_1 + \alpha \cos(\theta))$$

Here θ is the cosine of the angle between $\mathbf{x}$ and the space spanned by the two singular vectors. α is a control parameter that determines how close we want the vector to be to the subspace. When α is large, $\mathbf{x}$ must be close to the subspace. When α is zero, $\mathbf{x}$ can be anywhere.

Finding the minimum 1-norm vector in the space spanned by a set of vectors can be formulated as a linear programming problem. To find the minimum 1-norm vector in V , write $V\mathbf{x} = \mathbf{y}$ where we want to solve for both $\mathbf{x}$ and $\mathbf{y}$. Note that the format is different from the usual format for a set of linear equations $A\mathbf{x} = \mathbf{b}$ where $\mathbf{b}$ is a known vector.

Finding the minimum 1-norm vector looks like a nonlinear problem.

$$\min |\mathbf{y}|_1 \text{ subject to } V\mathbf{x} = \mathbf{y}$$

To remove the absolute value sign, write $\mathbf{y} = \mathbf{y}_1 - \mathbf{y}_2$ with $\mathbf{y}_1 \geq 0$ and $\mathbf{y}_2 \geq 0$. Then solve

$$\min \left(\sum_{i=1}^n y_{1i} + \sum_{i=1}^n y_{2i} \right) \text{ subject to } V\mathbf{x} = \mathbf{y}, \mathbf{y}_1 \geq 0, \text{ and } \mathbf{y}_2 \geq 0.$$

Write $V\mathbf{x} = \mathbf{y}_1 - \mathbf{y}_2$ as $V\mathbf{x} - \mathbf{y}_1 + \mathbf{y}_2 = 0$. then we have the linear equations in a format we are accustomed to.

$$[V, -I, I] \begin{pmatrix} \mathbf{x} \\ \mathbf{y}_1 \\ \mathbf{y}_2 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}$$

This is a linear programming problem. The solution, however, happens to be $\mathbf{x} = 0$, $\mathbf{y}_1 = 0$, and $\mathbf{y}_2 = 0$. To resolve this, add the equation $y_{1i} = 1$ to get a community containing the vertex i .

Often we are looking for communities of 50 or 100 vertices in graphs with hundreds of million of vertices. We want a method to find such communities in time proportional to the size of the community and not the size of the entire graph. Here spectral clustering can be used but instead of calculating singular vectors of the entire graph, we do something else. Consider a random walk on a graph. If we walk long enough the probability distribution converges to the first singular vector. However, if we take only a few steps from a

start vertex or small group of vertices that we believe define a cluster, the probability will distribute over the cluster with some of the probability leaking out to the remainder of the graph. To get the early convergence of several vectors which would ultimately converge to the first few singular vectors, we take a subspace $[\mathbf{x}, A\mathbf{x}, A^2\mathbf{x}, A^3\mathbf{x}]$ and propagate the subspace. At each iteration we find an orthonormal basis and then multiply each basis vector by A . We then take the resulting basis vectors after a few steps, say five, and find a minimum 1-norm vector in the subspace.

Spectral clustering applied to data

Consider n data points arranged as the rows of an $n \times d$ matrix A . These are to be partitioned into k clusters where k is much smaller than n or d . Finding the best k -means clustering of the points is known to be NP-hard. However, there are efficient algorithms that find a k -clustering within a factor of two of the best. We will see that singular value decomposition together with this type of approximate k -means clustering is very useful.

Spectral Clustering of the data into k clusters. The method consists of the following steps:

1. Find the top k right singular vectors of the data matrix A .
2. Project each row of A into the space spanned by these singular vectors to obtain a $n \times d$ matrix $\bar{A}$.
3. Apply an algorithm to find an approximately optimal k -clustering of $\bar{A}$.

It is important to note that the projected points are being clustered, i.e., the rows of $\bar{A}$ and not A itself. Projection offers the obvious advantage of decreasing the dimension of the problem from d to k , making it easier to cluster. The more important advantage of projecting is that it yields cluster centers closer to the true centers than clustering A . This is not so obvious and we demonstrate it here. The formal statement is contained in Theorem 8.7.

We will see how to use the fact that spectral clustering finds centers close to the true centers to get an actual clustering close to the true clustering. But this makes sense only if there is no ambiguity about what the true clustering is. We will develop a notion of a proper clustering that says the clusters are distinct enough so as not to be confused with each other. We will then show that if there is a proper clustering, spectral clustering will find a clustering close to the proper clustering. This is proved in Theorem 8.9.

Consider a spherical Gaussian F in $\mathbf{R}^d$ with mean $\boldsymbol{\mu}$ and variance one in every direction. As we saw in Chapter 2, for a point $\mathbf{x}$ picked according to F , $|\mathbf{x} - \boldsymbol{\mu}|^2$ is likely to be about d . So an approximate k -means clustering algorithm may come up with cluster centers $\boldsymbol{\mu}'$ with $|\boldsymbol{\mu}' - \boldsymbol{\mu}|^2 \leq d$ and still have $|\mathbf{x} - \boldsymbol{\mu}'|^2 \leq O(d)$.

Now, consider a mixture of k spherical Gaussians, each of variance one in every direction. We saw in Chapter 4 that the space spanned by the top k singular vectors contains the means of the k Gaussians. Project all data points on to this space. The densities are still Gaussian in the projection with variance again one in every direction. The mean squared distance of projected data points to the projected mean of the respective densities is $O(k)$ and if we find an approximately best k -means clustering, the cluster centers will be at distance squared at most $O(k)$ from the true means, not $O(d)$. In this example, we assumed that the data points were stochastically generated from a mixture of Gaussians. We show in what follows that this is not necessary. Indeed, we show that the intuitive argument here also holds for any arbitrary set of data points. But first, we have to define an analog of variance for a general set of data points. This is simple. It is just the average squared distance from the cluster center instead of the average distance squared to the mean of the probability density. Now, for spherical Gaussians, the squared distance in every direction is the same, but in general, they are not and we will take the maximum over all directions.

Represent a k -clustering by a $n \times d$ matrix C with each row of C being the cluster center of the cluster the corresponding row of A belongs to. Note that C has only k distinct rows. Define the *variance* of C , denoted $\sigma^2(C)$, by

$$\sigma^2(C) = \max_{\substack{\mathbf{v} \\ |\mathbf{v}|=1}} \frac{1}{n} |(A - C)\mathbf{v}|^2,$$

which is simply the maximum, in any direction $\mathbf{v}$, of the mean-squared distance of a data point from its cluster center. This is also $\frac{1}{n} \|A - C\|_2^2$.

If we had a stochastic model of data, as for example a mixture of Gaussians generating the data, then there is a true clustering and it is desirable that our algorithm find this clustering or at least come close. In general, we do not assume that there is a stochastic model and so there is no true clustering. Nevertheless, we will be able to show that spectral clustering does nearly as well as any clustering C . Namely, for most data points, the cluster centers found by spectral clustering will be at distance at most $O\left(\sqrt{k} \sigma(C)\right)$ of the cluster centers in C . The reader should think about the question: How is it that the one clustering found by the algorithm can do this for every possible clustering? For the theorem below, recall the notation that $\mathbf{a}_i$, $\mathbf{c}_i$, and $\mathbf{c}'_i$ are respectively the i^{th} row of A , C , and C' .

First, we need two technical lemmas.

Lemma 8.5 *For any two vectors $\mathbf{u}$ and $\mathbf{v}$,*

$$|\mathbf{u} + \mathbf{v}|^2 \geq \frac{1}{2} |\mathbf{u}|^2 - |\mathbf{v}|^2.$$

Proof:

$$\begin{aligned} |\mathbf{u} + \mathbf{v}|^2 &= (\mathbf{u} + \mathbf{v}) \cdot (\mathbf{u} + \mathbf{v}) = |\mathbf{u}|^2 + |\mathbf{v}|^2 + 2\mathbf{u} \cdot \mathbf{v} \\ &\geq |\mathbf{u}|^2 + |\mathbf{v}|^2 - 2|\mathbf{u}||\mathbf{v}| = (|\mathbf{u}| - |\mathbf{v}|)^2. \end{aligned}$$

From this and the fact that for any two real numbers a and b ,

$$(a - b)^2 \geq (a - b)^2 - \left(\frac{1}{\sqrt{2}}a - \sqrt{2}b \right)^2 = \frac{1}{2}a^2 - b^2,$$

the claim follows. ■

Lemma 8.6 *Suppose A is an $n \times d$ matrix and $\bar{A}$ is the projection of the rows of A onto the subspace spanned by the top k singular vectors of A . Then for any matrix C of rank at most k ,*

$$\|\bar{A} - C\|_F^2 \leq 8k\|A - C\|_2^2.$$

Proof: Since the rank of $\bar{A} - C$ is at most the sum of the ranks of $\bar{A}$ and C , which is most $2k$,

$$\|\bar{A} - C\|_F^2 \leq 2k\|\bar{A} - C\|_2^2 \tag{8.1}$$

by Lemma 4.2. Now

$$\|\bar{A} - C\|_2 \leq \|\bar{A} - A\|_2 + \|A - C\|_2 \leq 2\|A - C\|_2,$$

the last inequality since $\bar{A}$ is the best rank k approximation for the spectral norm and C has rank at most k . Combining this with (8.1), the lemma follows. ■

Theorem 8.7 *Suppose A is a $n \times d$ data matrix and C is any clustering of A and suppose C' (also a $n \times d$ matrix with k distinct rows) is the clustering of $\bar{A}$ found by the spectral clustering algorithm. For all but εn of the data points, we have $|\mathbf{c}_i - \mathbf{c}'_i|^2 < \frac{48k}{\varepsilon}\sigma^2(C)$.*

Proof: Let $\Delta = \frac{48k}{\varepsilon}\sigma^2(C)$ and $B = \{i \mid |\mathbf{c}_i - \mathbf{c}'_i|^2 \geq \Delta\}$ be the bad set of i . We must show that B has at most εn elements.

$$\begin{aligned} \sum_{i \in B} |\bar{\mathbf{a}}_i - \mathbf{c}'_i|^2 &= \sum_{i \in B} |(c_i - c'_i) + (\bar{a}_i - ci)|^2 \\ &\geq \frac{1}{2} \sum_{i \in B} |\mathbf{c}_i - \mathbf{c}'_i|^2 - \sum_{i \in B} |\bar{\mathbf{a}}_i - \mathbf{c}_i|^2 && \text{by the Lemma 8.5} \\ &\geq \frac{1}{2}|B|\Delta - \sum_{i=1}^n |\bar{\mathbf{a}}_i - \mathbf{c}_i|^2 = \frac{1}{2}|B|\Delta - \|\bar{A} - C\|_F^2. \end{aligned}$$

On the other hand,

$$\sum_{i \in B} |\bar{\mathbf{a}}_i - \mathbf{c}'_i|^2 \leq \sum_{i=1}^n |\bar{\mathbf{a}}_i - \mathbf{c}'_i|^2 \leq 2 \sum_{i=1}^n |\bar{\mathbf{a}}_i - \mathbf{c}_i|^2 = 2\|\bar{A} - C\|_F^2,$$

since, C' is within a factor of two of being the best k -means clustering of the projected data matrix $\bar{A}$ implies that if we took C as a clustering of $\bar{A}$, then, it is at most a factor of two better than C' . Combining,

$$3\|\bar{A} - C\|_F^2 \geq \frac{1}{2}|B|\Delta$$

which implies

$$|B| \leq \frac{6\varepsilon\|\bar{A} - C\|_F^2}{48k\sigma^2(C)}.$$

From Lemma 8.6, $\|\bar{A} - C\|_F^2 \leq 8k\|A - C\|_2^2 = 8kn\sigma^2(C)$. Plugging this in, the theorem follows. $\blacksquare$

We need the following lemma which asserts another property of spectral clustering, namely that the clustering it finds has σ which is within a factor of $5\sqrt{k}$ of the best possible σ for any clustering.

Lemma 8.8 *Let C^* be the k -clustering with the minimum σ among all k -clusterings of the data. For the clustering C' found by spectral clustering, we have*

$$\sigma(C') \leq 5\sqrt{k}\sigma(C^*).$$

Proof:

$$\begin{aligned} \|A - C'\|_2 &\leq \|A - \bar{A}\|_2 + \|\bar{A} - C'\|_2 \leq \|A - C^*\|_2 + \|\bar{A} - C'\|_F \\ &\leq \sqrt{n}\sigma(C^*) + \sqrt{2}\|\bar{A} - C^*\|_F \leq \sqrt{n}\sigma(C^*) + 4\sqrt{kn}\sigma(C^*), \text{ by Lemma 8.6.} \end{aligned}$$

For the second inequality, we used the fact that since $\bar{A}$ is the best rank k approximation to A in spectral norm, $\|A - C^*\|_2 \geq \|A - \bar{A}\|_2$ and for the third inequality, we used the fact that C' is within a factor of two of the optimal k -means clustering of $\bar{A}$ and in particular, the clustering C^* is not better by a factor of more than two. Now the lemma follows. $\blacksquare$

Now we show that we can use the fact that spectral clustering finds cluster centers close to the true centers to find approximately the true clustering. This makes sense only if there is no ambiguity about what the true clustering is. A necessary condition for a clustering to be unambiguous is that the clusters must be distinct or spatially well-separated. Otherwise, points could be put into either of two nearby clusters without changing the k -means objective function much. We make this more precise with the following definition:

Definition 8.1 *A clustering C^* is said to be proper if*

1. $\sigma(C^*)$ is least among all k -clustering of the data, and
2. the centers of any two clusters in C^* are separated by a distance of at least $70k^2\sigma(C^*)/\sqrt{\varepsilon}$.

■

Here, ε is any positive real number. Why do we choose this definition of a proper clustering? For the case of spherical Gaussians, the separation required here corresponds to the means of different Gaussians being a constant number of standard deviations apart. A different definition might have insisted on the clusters being distinct enough so that even an approximately best k -means clustering in $\mathbf{R}^d$, would give approximately the true clustering. Since for a spherical Gaussian with variance one in every direction, data points are about $\sqrt{d}$ away from the center, this would intuitively require the means of two such Gaussians involved in a mixture to be at least $\Omega(\sqrt{d})$ apart, which is a stronger requirement than that of being proper. To be proper, a separation of only $\Omega(1)$ is required.

We modify the spectral clustering algorithm by adding a *merge step* at the end.

Merge Step Let C' be the clustering found by spectral clustering. Repeatedly merge any two clusters with cluster centers separated by a distance of at most $14\sqrt{k}\sigma(C')/\sqrt{\varepsilon}$.

Theorem 8.9 *Suppose there is a proper clustering C^* of data points. Then, spectral clustering followed by merge-step produces a clustering $C^{(0)}$ with the property that by reclustering at most εn points, we can get from $C^{(0)}$ to C^* .*

Proof: Let C' be the clustering produced by spectral clustering before the merge step is executed. Let $\Delta = 49k\sigma^2(C')/\varepsilon$. Define $B = \{i : |\mathbf{c}'_i - \mathbf{c}^*_i|^2 > \Delta\}$. Let S be one particular cluster in C^* . For any $i, j \in S \setminus B$, we have $|\mathbf{c}'_i - \mathbf{c}^*_i| \leq \sqrt{\Delta}$ and $|\mathbf{c}'_j - \mathbf{c}^*_j| \leq \sqrt{\Delta}$. Since $\mathbf{c}^*_i = \mathbf{c}^*_j$, i and j will be in one cluster after the merge step. Now if S and T are two different clusters in C^* , by the definition of proper, for $i \in S \setminus B$ and $j \in T \setminus B$, we have

$$|\mathbf{c}^*_i - \mathbf{c}^*_j| \geq 70k^2\sigma(C^*)/\sqrt{\varepsilon} \geq 14k^{3/2}\sigma(C')/\sqrt{\varepsilon} \geq 2k\sqrt{\Delta},$$

by Lemma 8.8. So $|\mathbf{c}'_i - \mathbf{c}'_j| \geq 2(k-1)\sqrt{\Delta}$ by the definition of B and the merge step (even when repeated $k-1$ times) will not merge i and j into one cluster. Thus, for all $i, j \notin B$, we have that i and j belong to the same cluster in C^* if and only if they belong to the same cluster in $C^{(0)}$. Thus, by reclustering at most $|B|$ points, we can get from $C^{(0)}$ to C^* . ■

8.5 Recursive Clustering Based on Sparse Cuts

Suppose we are given an undirected, connected graph $G(V, E)$ in which an edge indicates the end point vertices are similar. Recursive clustering starts with all vertices in one cluster and recursively splits a cluster into two parts whenever there are not too many edges from one part to the other part of the cluster. For this technique to be effective it is important that the data has an hierarchical clustering. Consider what would happen if one used recursive clustering to find communities of students at an institution hoping that one of the clusters might be computer science students. At the first level one

might get four clusters corresponding to freshman, sophomores, juniors, and seniors. At the next level one might get clusters that were majors partitioned into year rather than majors. Another problem would occur if the real top level clusters were overlapping. If one is clustering journals articles, the top level might be mathematics, physics, chemistry, etc. However, there are papers that are related to both mathematics and physics. Such a paper would be put in one cluster or the other and the community that the paper really belonged in would be split and thus never found at lower levels in the clustering.

Formally, for two disjoint sets S and T of vertices, define

$$\Phi(S, T) = \frac{\text{number of edges from } S \text{ to } T}{\text{total number of edges incident to } S \text{ in } G}.$$

$\Phi(S, T)$ measures the relative strength of similarities between S and T . Let $d(i)$ be the degree of vertex i and for $d(S) = \sum_{i \in S} d(i)$. Let m be the total number of edges. The following algorithm cuts only a small fraction of the edges, yet ensures that each cluster is consistent, namely no subset of it has low similarity to the rest of the cluster.

Recursive Clustering Algorithm

If a current cluster W has a subset S with $d(S) \leq \frac{1}{2}d(W)$ and $\Phi(S, W-S) \leq \varepsilon$, then split W into two clusters: $S, W-S$. Repeat until no such split is possible.

Theorem 8.10 *At termination of the above algorithm, the total number of edges between vertices in different clusters is at most $O(\varepsilon m \ln n)$.*

Proof: Each edge between two different clusters at the end was “cut up” at some stage by the algorithm. We will “charge” edge cuts to vertices and bound the total charge. When the algorithm partitions a cluster W into S and $W-S$ with $d(S) \leq (1/2)d(W)$, each $k \in S$ is charged $\frac{d(k)}{d(W)}$ times the number of edges being cut. Since $\Phi(S, W-S) \leq \varepsilon$, the charge added to each $k \in W$ is at most $\varepsilon d(k)$. A vertex is charged only when it is in the smaller part ($d(S) \leq d(W)/2$) of the cut. So between any two times it is charged, $d(W)$ is reduced by a factor of at least two and so a vertex can be charged at most $\log_2 m \leq O(\ln n)$ times, proving the theorem. ■

To implement the algorithm, we have to compute $\text{Min}_{S \subseteq W} \Phi(S, W-S)$, an NP-hard problem. So the theorem cannot be implemented right away. Luckily, eigenvalues and eigenvectors, which can be computed fast, give an approximate answer. The connection between eigenvalues and sparsity, known as Cheeger’s inequality, is deep with applications to Markov chains among others. We do not discuss this here.

8.6 Kernel Methods

The clustering methods discussed so far work well only when the data satisfy certain conditions. For example, in any distance-based measure like k -means or k -center, once

the cluster centers are fixed, the Vornoi diagram of the cluster centers determines which cluster each data point belongs to. Cells of the Vornoi diagram are determined by hyperplane bisectors of line segments joining pairs of centers. This implies that clusters are linearly separable.

Such criteria cannot separate clusters that are not linearly separable in the input space. In the chapter on learning had many examples that were not linearly separable in the original space, but were linearly separable when mapped to a higher dimensional space using a nonlinear function called a kernel. An analogous technique can be used in the case of clustering, but with two differences.

1. There may be any number k of clusters, whereas in learning, there were just two classes, the positive and negative examples.
2. There is unlabelled data, i.e., we are not given which cluster each data point belongs to, whereas in the case of learning each data point was labeled. The clustering situation is sometimes called *unsupervised* whereas the labeled learning situation is called *supervised*, the reason being, one imagines a supervisor, human judgement, supplying the labels.

These two differences do not prevent the application of kernel methods to clustering. Indeed, here too, one could first embed the data in a different space using the Gaussian or other kernel and then run k -means in the embedded space. Again, one need not write down the whole embedding explicitly. In the learning setting, since there were only two classes with a linear separator, we were able to write a convex program to find the normal of the separator. When there are k classes, there could be as many as $\binom{k}{2}$ hyperplanes separating pairs of classes, so the computational problem is harder and there is no simple convex program to solve the problem. However, we can still run the k -means algorithm in the embedded space. The centroid of a cluster is kept as the average of the data points in the cluster. Recall that we only know dot products, not distances in the higher dimensional space, but we can use the relation $\|\mathbf{x} - \mathbf{y}\|^2 = \mathbf{x} \cdot \mathbf{x} + \mathbf{y} \cdot \mathbf{y} - 2\mathbf{x} \cdot \mathbf{y}$ to go from dot products to distances.

There are situations in which high-dimensional data points lie in a lower dimensional manifold. In such situations, a Gaussian kernel is useful. Say we are given a set of n points $S = \{\mathbf{s}_1, \mathbf{s}_2, \dots, \mathbf{s}_n\}$ in $\mathbf{R}^d$ that we wish to cluster into k subsets. The Gaussian kernel uses an affinity measure that emphasizes closeness of points and drops off exponentially as the points get farther apart. We define the affinity between points i and j by

$$a_{ij} = \begin{cases} e^{-\frac{1}{2\sigma^2} \|\mathbf{s}_i - \mathbf{s}_j\|^2} & i \neq j \\ 0 & i = j. \end{cases}$$

The affinity matrix gives a closeness measure for points. The measure drops off exponentially fast with distance and thus favors close points. Points farther apart have their

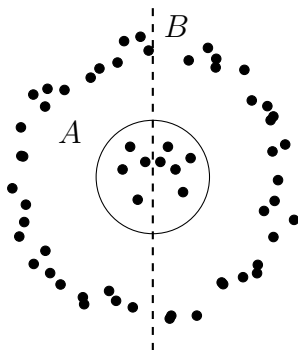


Figure 8.2: Example where 2-median clustering B is not natural.

closeness shrink to zero. We give two examples to illustrate briefly the use of Gaussian kernels. The first example is similar to Figure 8.2 of points on two concentric annuli. Suppose the annuli are close together, i.e., the distance between them is $\delta \ll 1$. Even if we used similarity between objects, rather than say the k -median criterion, it is not clear that we will get the right clusters; namely two separate circles. Instead, suppose the circles are sampled at a rate so that adjacent samples are separated by a distance $\varepsilon \ll \delta$. Define a Gaussian kernel with variance ε^2 . Then, if sample $\mathbf{s}_1$ is on Circle 1 and sample $\mathbf{s}_2$ is on Circle 2, $e^{-|\mathbf{s}_1 - \mathbf{s}_2|^2 / 2\varepsilon^2} \ll 1$, so they are very likely to be put in separate clusters as desired.

Our second example has three curves. Suppose two points from two different curves are never closer than δ . If we sample at a high enough rate, every sample will have many other samples from the same curve close to it giving high similarity according to the Gaussian kernel, but no two samples from different curves will have high similarity. This example can be generalized to a situation where the points lie on different “sheets” or low dimensional manifolds.

Two points near each other on the same circle will have a high affinity value, i.e., they will be close together in this metric. For the right sigma value, the two closest points, one from each circle, will be infinitely far apart. Thus, the affinity matrix is a band matrix, consisting of two blocks of data.

8.7 Agglomerative Clustering

Agglomerative clustering is the opposite of recursive clustering. It starts with each point in a separate cluster and then repeatedly merges the two closest clusters into one. There are various criteria to determine which two clusters are merged at any point. They are based on first defining a distance between two clusters in terms of the distance between points. Four of these possibilities are listed below.

1. Nearest neighbor - the distance between clusters C_i and C_j is the distance between

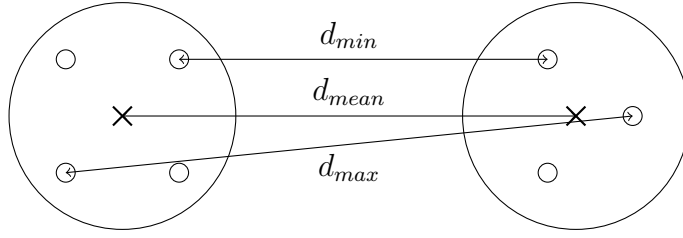


Figure 8.3: Illustration of minimum, maximum and mean distances for a set of points.

the points in C_i and C_j that are closest.

$$d_{\min}(C_i, C_j) = \min_{\substack{\mathbf{x} \in C_i \\ \mathbf{y} \in C_j}} |\mathbf{x} - \mathbf{y}|$$

This measure basically builds the minimal cost spanning tree.

2. Farthest neighbor - the distance between clusters C_i and C_j is the distance between the points in C_i and C_j that are farthest apart.

$$d_{\max}(C_i, C_j) = \max_{\substack{\mathbf{x} \in C_i \\ \mathbf{y} \in C_j}} |\mathbf{x} - \mathbf{y}|$$

3. Mean - the distance between two clusters is the distance between the centroids of the clusters.
4. Average - the distance between two clusters is the average distance between points in the two clusters.

Agglomerative clustering in moderate and high dimensions often gives rise to a very unbalanced tree. This section gives some insight into the cause of this phenomenon. We begin by considering agglomerative clustering of 1-dimensional data. The 1-dimensional case gives rise to a balanced tree.

Consider agglomerative clustering using nearest neighbor of n points on a line. Assume the distances between adjacent points are independent random variables and measure the distance between clusters by the distance of the nearest neighbors. Here it is easy to see that each cluster is always an interval. In this case, any two adjacent intervals are equally likely to be merged. The last merge will occur when the largest distance between two adjacent points is encountered. This is equally likely to be the distance between point i and point $i + 1$ for any i , $1 \leq i < n$. The height of the final merge tree will be one more than the maximum of the heights of the two subtrees. Let $h(n)$ be the expected height of a merge tree with n leaves. Then

$$\begin{aligned}
h(n) &= 1 + \frac{1}{n} \sum_{i=1}^n \max \{h(i), h(n-i)\} \\
&= 1 + \frac{2}{n} \sum_{i=\frac{n}{2}+1}^n h(i) \\
&= 1 + \frac{2}{n} \sum_{i=\frac{n}{2}+1}^{\frac{3}{4}n} h(i) + \frac{2}{n} \sum_{i=\frac{3}{4}n+1}^n h(i)
\end{aligned}$$

Since $h(i)$ is monotonic, for $\frac{n}{2} < i \leq \frac{3}{4}n$, bound $h(i)$ by $h(\frac{3}{4}n)$ and for $\frac{3}{4}n < i \leq n$, bound $h(i)$ by $h(n)$. Thus,

$$\begin{aligned}
h(n) &\leq 1 + \frac{2}{n} \frac{n}{4} h\left(\frac{3n}{4}\right) + \frac{2}{n} \frac{n}{4} h(n) \\
&\leq 1 + \frac{1}{2} h\left(\frac{3n}{4}\right) + \frac{1}{2} h(n).
\end{aligned}$$

This recurrence has a solution $h(n) \leq b \log n$ for sufficiently large b . Thus, the merge tree has no long path and is bushy.

If the n points are in high dimension rather than constrained to a line, then the distance between any two points, rather than two adjacent points, can be the smallest distance. One can think of edges being added one at a time to the data to form a spanning tree. Only now we have an arbitrary tree rather than a straight line. The order in which the edges are added corresponds to the order in which the connected components are merged by the agglomerative algorithm. Two extreme cases of the spanning tree for the set of points are the straight line which gives a bushy agglomerative tree or a star which gives a skinny agglomerative tree of height n . Note there are two trees involved here, the spanning tree and the agglomerative tree.

The question is what is the shape of the spanning tree? If distance between components is nearest neighbor, the probability of an edge between two components is proportional to the size of the components. Thus, once a large component forms it will swallow up small components giving a more star like spanning tree and hence a tall skinny agglomerative tree. Notice the similarity to the $G(n, p)$ problem.

If we defined distance between two clusters to be the maximum distance between any two points in the clusters and merge the two clusters that are the smallest distance apart, then we are more likely to get a bushy spanning tree and a skinny agglomerative tree. If all distances between points are independent and we have two clusters of size k and a singleton, the maximum distance between the points in the two clusters of size k is likely

to give a larger distance than the maximum between the singleton and the k points in a cluster. Thus, the singleton will likely merge into one of the clusters before the two clusters will merge and in general small clusters will combine before larger ones, resulting in a bushy spanning tree and a bushy agglomerative tree.

8.8 Dense Submatrices and Communities

Represent n data points in d -space by the rows of an $n \times d$ matrix A . Assume that A has all nonnegative entries. Examples to keep in mind for this section are the document-term matrix and the customer-product matrix. We address the question of how to define and find efficiently a coherent large subset of rows. To this end, the matrix A can be represented by a bipartite graph. One side has a vertex for each row and the other side a vertex for each column. Between the vertex for row i and the vertex for column j , there is an edge with weight a_{ij} .

We want a subset S of row vertices and a subset T of column vertices so that

$$A(S, T) = \sum_{i \in S, j \in T} a_{ij}$$

is high. This simple definition is not good since $A(S, T)$ will be maximized by taking all rows and columns. We need a balancing criterion that ensures that $A(S, T)$ is high relative to the sizes of S and T . One possibility is to maximize $\frac{A(S, T)}{|S||T|}$. This is not a good measure either, since it is maximized by the single edge of highest weight. The definition we use is the following. Let A be a matrix with nonnegative entries. For a subset S of rows and a subset T of columns, the *density* $d(S, T)$ of S and T is $d(S, T) = \frac{A(S, T)}{\sqrt{|S||T|}}$. The *density* $d(A)$ of A is defined as the maximum value of $d(S, T)$ over all subsets of rows and columns. This definition applies to bipartite as well as non bipartite graphs.

One important case is when A 's rows and columns both represent the same set and a_{ij} is the similarity between object i and object j . Here $d(S, S) = \frac{A(S, S)}{|S|}$. If A is an $n \times n$ 0-1 matrix, it can be thought of as the adjacency matrix of an undirected graph, and $d(S, S)$ is the average degree of a vertex in S . The subgraph of maximum average degree in a graph can be found exactly by network flow techniques, as we will show in the next section. We do not know an efficient (polynomial-time) algorithm for finding $d(A)$ exactly in general. However, we show that $d(A)$ is within a $O(\log^2 n)$ factor of the top singular value of A assuming $|a_{ij}| \leq 1$ for all i and j . This is a theoretical result. The gap may be much less than $O(\log^2 n)$ for many problems, making the singular value and singular vector quite useful. Also, S and T with $d(S, T) \geq \Omega(d(A)/\log^2 n)$ can be found algorithmically.

Theorem 8.11 *Let A be an $n \times d$ matrix with entries between 0 and 1. Then*

$$\sigma_1(A) \geq d(A) \geq \frac{\sigma_1(A)}{4 \log n \log d}.$$

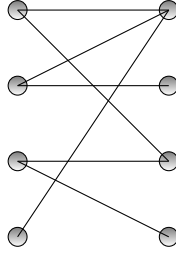


Figure 8.4: Example of a bipartite graph.

Furthermore, subsets S and T satisfying $d(S, T) \geq \frac{\sigma_1(A)}{4 \log n \log d}$ may be found from the top singular vector of A .

Proof: Let S and T be the subsets of rows and columns that achieve $d(A) = d(S, T)$. Consider an n -vector $\mathbf{u}$ which is $\frac{1}{\sqrt{|S|}}$ on S and 0 elsewhere and a d -vector $\mathbf{v}$ which is $\frac{1}{\sqrt{|T|}}$ on T and 0 elsewhere. Then,

$$\sigma_1(A) \geq \mathbf{u}^T A \mathbf{v} = \sum_{ij} u_i v_j a_{ij} = d(S, T) = d(A)$$

establishing the first inequality.

To prove the second inequality, express $\sigma_1(A)$ in terms of the first left and right singular vectors $\mathbf{x}$ and $\mathbf{y}$.

$$\sigma_1(A) = \mathbf{x}^T A \mathbf{y} = \sum_{i,j} x_i a_{ij} y_j, \quad |\mathbf{x}| = |\mathbf{y}| = 1.$$

Since the entries of A are nonnegative, the components of the first left and right singular vectors must all be positive, that is, $x_i \geq 0$ and $y_j \geq 0$ for all i and j . To bound $\sum_{i,j} x_i a_{ij} y_j$, break the summation into $O(\log n \log d)$ parts. Each part corresponds to a given α and β and consists of all i such that $\alpha \leq x_i < 2\alpha$ and all j such that $\beta \leq y_j < 2\beta$. The $\log n \log d$ parts are defined by breaking the rows into $\log n$ blocks with α equal to $\frac{1}{2\sqrt{n}}, \frac{1}{\sqrt{n}}, 2\frac{1}{\sqrt{n}}, 4\frac{1}{\sqrt{n}}, \dots, 1$ and by breaking the columns into $\log d$ blocks with β equal to $\frac{1}{2\sqrt{d}}, \frac{1}{\sqrt{d}}, \frac{2}{\sqrt{d}}, \frac{4}{\sqrt{d}}, \dots, 1$. The i such that $x_i < \frac{1}{2\sqrt{n}}$ and the j such that $y_j < \frac{1}{2\sqrt{d}}$ will be ignored at a loss of at most $\frac{1}{4}\sigma_1(A)$. [Exercise (8.28) proves the loss is at most this amount.]

Since $\sum_i x_i^2 = 1$, the set $S = \{i | \alpha \leq x_i < 2\alpha\}$ has $|S| \leq \frac{1}{\alpha^2}$ and similarly,

$T = \{j | \beta \leq y_j \leq 2\beta\}$ has $|T| \leq \frac{1}{\beta^2}$. Thus

$$\begin{aligned} \sum_{\alpha \leq x_i \leq 2\alpha} \sum_{\beta \leq y_j \leq 2\beta} x_i y_j a_{ij} &\leq 4\alpha\beta A(S, T) \\ &\leq 4\alpha\beta d(S, T) \sqrt{|S||T|} \\ &\leq 4d(S, T) \\ &\leq 4d(A). \end{aligned}$$

From this it follows that

$$\sigma_1(A) \leq 4d(A) \log n \log d$$

or

$$d(A) \geq \frac{\sigma_1(A)}{4 \log n \log d}$$

proving the second inequality.

It is also clear that for each of the values of (α, β) , we can compute $A(S, T)$ and $d(S, T)$ as above and taking the best of these $d(S, T)$'s gives us an algorithm as claimed in the Theorem. ■

Note that in many cases, the nonzero values of x_i and y_j (after zeroing out the low entries) will only go from $\frac{1}{2} \frac{1}{\sqrt{n}}$ to $\frac{c}{\sqrt{n}}$ for x_i and $\frac{1}{2} \frac{1}{\sqrt{d}}$ to $\frac{c}{\sqrt{d}}$ for y_j , since the singular vectors are likely to be balanced given that a_{ij} are all between 0 and 1. In this case, there will be $O(1)$ groups only and the log factors disappear.

Another measure of density is based on similarities. Recall that the similarity between objects represented by vectors (rows of A) is defined by their dot products. Thus, similarities are entries of the matrix AA^T . Define the average cohesion $f(S)$ of a set S of rows of A to be the sum of all pairwise dot products of rows in S divided by $|S|$. The average cohesion of A is the maximum over all subsets of rows of the average cohesion of the subset.

Since the singular values of AA^T are squares of singular values of A , we expect $f(A)$ to be related to $\sigma_1(A)^2$ and $d(A)^2$. Indeed it is. We state the following without proof.

Lemma 8.12 $d(A)^2 \leq f(A) \leq d(A) \log n$. Also, $\sigma_1(A)^2 \geq f(A) \geq \frac{c\sigma_1(A)^2}{\log n}$.

$f(A)$ can be found exactly using flow techniques as we will see later.

In this section, we described how to find a large global community. There is another question, that of finding a small local community including a given vertex. We will visit this question in Section 8.10.

8.9 Flow Methods

Here we consider dense induced subgraphs of a graph. An induced subgraph of a graph consisting of a subset of the vertices of the graph along with all edges of the graph that connect pairs of vertices in the subset of vertices. We show that finding an induced subgraph with maximum average degree can be done by network flow techniques. This is simply maximizing density $d(S, S)$ of Section 8.8 over all subsets S of the graph. First consider the problem of finding a subset of vertices such that the induced subgraph has average degree at least λ for some parameter λ . Then do a binary search on the value of λ until the maximum λ for which there exists a subgraph with average degree at least λ is found.

Given a graph G in which one wants to find a dense subgraph, construct a directed graph H from the given graph and then carry out a flow computation on H . H has a node for each edge of the original graph, a node for each vertex of the original graph, plus two additional nodes s and t . There is a directed edge with cost one from s to each node corresponding to an edge of the original graph and a directed edge with infinite cost from each node corresponding to an edge of the original graph to the two nodes corresponding to the vertices the edge connects. Finally, there is a directed edge with cost λ from each node corresponding to a vertex of the original graph to t .

Notice there are three types of cut sets of the directed graph that have finite capacity. The first cuts all arcs from the source. It has capacity e , the number of edges of the original graph. The second cuts all edges into the sink. It has capacity λv , where v is the number of vertices of the original graph. The third cuts some arcs from s and some arcs into t . It partitions the set of vertices and the set of edges of the original graph into two blocks. The first block contains the source node s , a subset of the edges e_s , and a subset of the vertices v_s defined by the subset of edges. The first block must contain both end points of each edge in e_s ; otherwise an infinite arc will be in the cut. The second block contains t and the remaining edges and vertices. The edges in this second block either connect vertices in the second block or have one endpoint in each block. The cut set will cut some infinite arcs from edges not in e_s coming into vertices in v_s . However, these arcs are directed from nodes in the block containing t to nodes in the block containing s . Note that any finite cost cut that leaves an edge node connected to s must cut the two related vertex nodes from t . Thus, there is a cut of cost $e - e_s + \lambda v_s$ where v_s and e_s are the vertices and edges of a subgraph. For this cut to be the minimal cut, the quantity $e - e_s + \lambda v_s$ must be minimal over all subsets of vertices of the original graph and the cost must be less than e and also less than λv .

If there is a subgraph with v_s vertices and e_s edges where the ratio $\frac{e_s}{v_s}$ is sufficiently large so that $\frac{e_s}{v_s} > \frac{e}{v}$, then for λ such that $\frac{e_s}{v_s} > \lambda > \frac{e}{v}$, $e_s - \lambda v_s > 0$ and $e - e_s + \lambda v_s < e$. Similarly $e < \lambda v$ and thus $e - e_s + \lambda v_s < \lambda v$. This implies that the cut $e - e_s + \lambda v_s$ is less than either e or λv and the flow algorithm will find a nontrivial cut and hence a proper

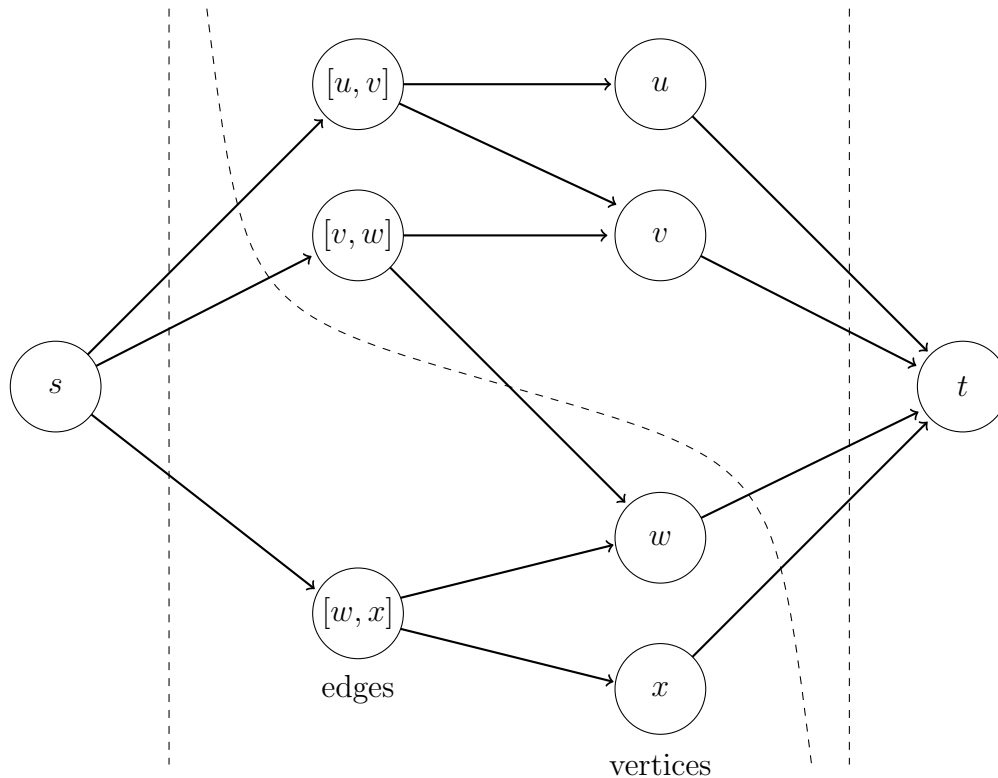


Figure 8.5: The directed graph H used by the flow technique to find a dense subgraph

subset. For different values of λ in the above range there maybe different nontrivial cuts.

Note that for a given density of edges, the number of edges grows as the square of the number of vertices and $\frac{e_s}{v_s}$ is less likely to exceed $\frac{e}{v}$ if v_s is small. Thus, the flow method works well in finding large subsets since it works with $\frac{e_s}{v_s}$. To find small communities one would need to use a method that worked with $\frac{e_s}{v_s^2}$ as the following example illustrates.

Example: Consider finding a dense subgraph of 1,000 vertices and 2,000 internal edges in a graph of 10^6 vertices and 6×10^6 edges. For concreteness, assume the graph was generated by the following process. First, a 1,000-vertex graph with 2,000 edges was generated as a random regular degree four graph. The 1,000-vertex graph was then augmented to have 10^6 vertices and edges were added at random until all vertices were of degree 12. Note that each vertex among the first 1,000 has four edges to other vertices among the first 1,000 and eight edges to other vertices. The graph on the 1,000 vertices is much denser than the whole graph in some sense. Although the subgraph induced by the 1,000 vertices has four edges per vertex and the full graph has twelve edges per vertex, the probability of two vertices of the 1,000 being connected by an edge is much higher than for the graph as a whole. The probability is given by the ratio of the actual number of edges connecting

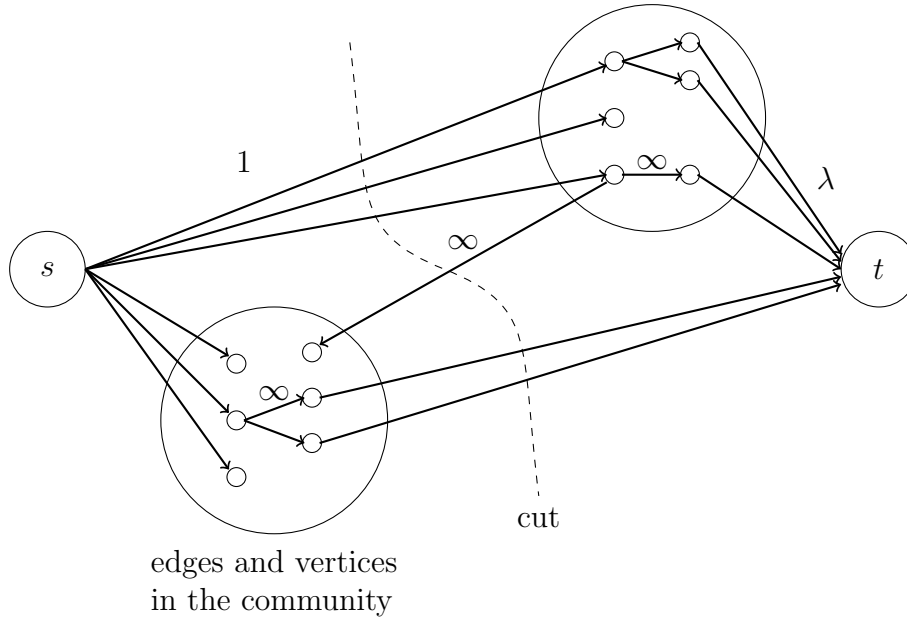


Figure 8.6: Cut in flow graph

vertices among the 1,000 to the number of possible edges if the vertices formed a complete graph. $\frac{A(S,S)}{|S|^2}$?

$$p = \frac{e}{\binom{v}{2}} = \frac{2e}{v(v-1)}$$

For the 1,000 vertices, this number is $p = \frac{2 \times 2,000}{1,000 \times 999} \cong 4 \times 10^{-3}$. For the entire graph this number is $p = \frac{2 \times 6 \times 10^6}{10^6 \times 10^6} = 12 \times 10^{-6}$. This difference in probability of two vertices being connected should allow us to find the dense subgraph. ■

In our example, the cut of all arcs out of s is of cost 6×10^6 , the total number of edges in the graph, and the cut of all arcs into t is of cost λ times the number of vertices or $\lambda \times 10^6$. A cut separating the 1,000 vertices and 2,000 edges would have cost $6 \times 10^6 - 2,000 + \lambda \times 1,000$. This cut cannot be the minimum cut for any value of λ since $\frac{e_s}{v_s} = 2$ and $\frac{e}{v} = 6$, hence $\frac{e_s}{v_s} < \frac{e}{v}$. The point is that to find the 1,000 vertices, we have to maximize $A(S,S)/|S|^2$ rather than $A(S,S)/|S|$. Note that $A(S,S)/|S|^2$ penalizes large $|S|$ much more and therefore can find the 1,000 node “dense” subgraph.

8.10 Finding a Local Cluster Without Examining the Whole Graph

If one wishes to find the community containing a vertex v in a large graph with say a billion vertices, one would like to find the community in time proportional to the size of

the community and independent of the size of the graph. Thus, we would like local methods that do not inspect the entire graph but only the neighborhood around the vertex v . We now give several such algorithms. Throughout this section, we assume the graph is undirected.

Breadth-First Search

The simplest method is to do a breadth first search starting at v . Clearly if there is a small connected component containing v , we will find it in time depending only on the size (number of edges) of the component. In a more subtle situation, each edge may have a weight that is the similarity between the two end points. If there is a small cluster C containing v , with each outgoing edge from C to $\bar{C}$ having weight less than some ε , C could clearly also be found by breadth-first search in time proportional to the size of C . However, in general, it is unlikely that the cluster will have such obvious telltale signs of its boundary and one needs more complex techniques, some of which we describe now.

By max flow

Given a vertex v in a directed graph, we want to find a small set S of vertices whose boundary (set of a few outgoing edges) is very small. Suppose we are looking for a set S whose boundary is of size at most b and whose cardinality is at most k . Clearly, if $\deg(v) < b$ then the problem is trivial, so assume $\deg(v) \geq b$.

Think of a flow problem where v is the source. Put a capacity of one on each edge of the graph. Create a new vertex that is the sink and add an edge of capacity α from each vertex of the original graph to the new sink vertex, where $\alpha = b/k$. If a community of size at most k with boundary at most b containing v exists, then there will be a cut separating v from the sink of size at most $k\alpha + b = 2b$, since the cut will have k edges from the community to the sink and b edges from the community to the remainder of the graph. Conversely, if there is a cut of size at most $2b$, then the community containing v has a boundary of size at most $2b$ and has at most $2k$ vertices since each vertex has an edge to the sink with capacity $\frac{b}{k}$. Thus, to come within a factor of two of the answer, all one needs to do is determine whether there is a cut of size at most $2b$. Since we know that the minimum size of any cut equals the maximum flow, it suffices to find the maximum flow. If the flow algorithm can do more than $2k$ flow augmentations, then the maximum flow and hence the minimum cut is of size more than $2b$. If not, the minimum cut is of size at most $2b$.

In executing the flow algorithm one finds an augmenting path from source to sink and augments the flow. Each time a new vertex not seen before is reached, there is an edge to the sink and the flow can be augmented by α directly on the path from v to the new vertex to the sink. So the amount of work done is a function of b and k , not the total number of vertices in the graph.

Sparsity and Local communities

In this part, we consider another definition of a local community. A local community in an undirected graph $G(V, E)$ is a subset of vertices with strong internal similarities and weak similarities to the outside. Using the same notation as in Section 8.5, we formalize this as follows:

Definition 8.2 *A subset S of vertices is a local community with parameter $\varepsilon > 0$ if it satisfies the following conditions:*

$$\Phi(S, \bar{S}) \leq \varepsilon^3 \quad (8.2)$$

$$\forall T \subseteq S, d(T) \leq \frac{1}{2}d(S), \quad \Phi(T, S \setminus T) \geq \varepsilon. \quad (8.3)$$

■

The first condition says that the connections of S to the outside $\bar{S}$ are weak. The second condition requires subsets of S of size as measured by $d(\cdot)$ less than $1/2$ of it to be strongly connected to the rest of S . Otherwise, S would not be one community, rather it would split into at least two. Note that for $\varepsilon \ll 1$, we have $\varepsilon^3 \ll \varepsilon$ and so the internal connections are required to be much stronger than the external ones. This is intuitively consistent with what we think of as a strong community. However, as opposed to Section 8.5, where, we spent time that grows as a function of $|V|$ since recursive clustering starts with the whole of V as one cluster, here, we will assume that $|S| \ll |U|$ and would like to find S in time which grows as a function of $|S|$, not $|V|$.

To accomplish this, we do a random walk on the graph starting with a vertex in S with transition probability matrix P (See Chapter 5) given by:

$$p_{ij} = \frac{1}{d_i} \quad \text{for } j \text{ adjacent to } i.$$

Recall from Chapter 5 that the Fundamental Theorem of Markov Chains proved that the long-term average probability vector converges to a stationary probability π given by

$$\pi_i = \frac{d(i)}{\sum_{j \in V} d(j)},$$

assuming G is connected.

If the Markov Chain is run for long enough, the probabilities will “spread” throughout V in proportion to the degrees. This is not desirable in this context. We would rather have it be essentially confined to S , our local community. Intuitively, since S ’s connection to $\bar{S}$ is at most ε^3 , if we run the Markov Chain for $O(1/\varepsilon^2)$ steps, we hope to have only ε probability of stepping into $\bar{S}$. Unfortunately, this is not valid. There may be a few

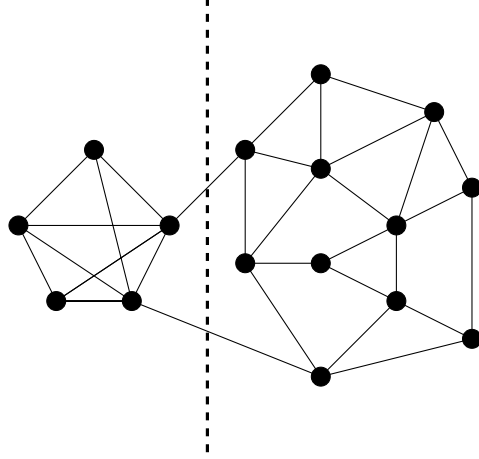


Figure 8.7: Example of a highly connected community of low sparsity.

boundary vertices in S that have strong connections to $\bar{S}$ and if we happen to start in one of them, we might step into $\bar{S}$ right away. All we can assert is that for most starting points in S , we do not step into $\bar{S}$ in $O(1/\varepsilon^2)$ steps. We now show this.

Lemma 8.13 *Suppose condition (8.2) is satisfied. Then there is a subset S_0 of S with $\pi(S_0) \geq \frac{3}{4}\pi(S)$ such that starting the Markov Chain at any i in S_0 and running it for $t_0 \in O(1/\varepsilon^2)$ steps, the probability that we ever step into $\bar{S}$ is at most $O(\varepsilon)$.*

Proof: For $i \in S$ and $t \in O(1/\varepsilon^2)$, let $f(i, t)$ be the probability that the Markov chain started at vertex i at time 0, walks from S to a vertex in $\bar{S}$ at time t . We would like to upper bound $f(i, t)$. While this is difficult, observe that had the walk started in the stationary distribution π , then it would remain in the stationary distribution and so the probability that it would step from a vertex in S to a vertex in $\bar{S}$ at time t is precisely $\sum_{j \in S} \sum_{k \in \bar{S}} \pi_j p_{jk} = \Phi(S)\pi(S)$. By linearity, this probability is $\sum_{i \in S} \pi_i f(i, t)$, so

$$\sum_{i \in S} \pi_i f(i, t) = \Phi(S)\pi(S) \leq \varepsilon^3 \pi(S)$$

Let f_i be the probability that when started at time 0 at i , we step at least once into $\bar{S}$ in the first $O(1/\varepsilon^2)$ steps of the Markov Chain. We get from above that:

$$\sum_{i \in S} \frac{\pi_i}{\pi(S)} f_i \leq O(\varepsilon).$$

But $\sum_{i \in S} \frac{\pi_i}{\pi(S)} f_i$ is the weighted average of f_i over $i \in S$ with weights $\pi_i/\pi(S)$ and so by Markov inequality, it follows that the weight of the set of i for which $f_i > c\varepsilon$ for a large constant c cannot exceed $1/4$ proving the lemma. ■

To discover S , besides not going out of S , we also need that the walk spreads through all, or at least most, of S which we now show.

Lemma 8.14 *Suppose conditions (8.2) and (8.3) are satisfied. Start the markov Chain in S_0 and run it for t_0 steps. Let S_1 be the set of all $i \in S$ for which the expected number of visits is at least $(3/4)\pi_i/\pi(S)$. Then, $\pi(S_1) \geq (3/4)\pi(S)$.*

Proof: For $i \in S$ and $t \leq t_0$, let

$$g_{it} = \text{Prob}(\text{Walk remains in } S \text{ and is at } i \text{ at time } t)$$

$$h_i = \frac{1}{t_0} \sum_{t=0}^{t_0-1} g_{it}.$$

From Lemma 8.13, we have

$$\sum_{i \in S} g_{it} = \text{Prob}(\text{Walk remains in } S) \geq 1 - O(\varepsilon) \implies \sum_{i \in S} h_i \geq 1 - O(\varepsilon). \quad (8.4)$$

Let $\tilde{P}$ be the transition probability matrix of a Markov chain with states S obtained from P by redirecting each transition (i, j) from a vertex $i \in S$ to a vertex $j \in \bar{S}$ to be a self-loop at i . More formally, $\tilde{p}_{jk} = p_{jk}$, for $j, k \in S$ and for all $i \in S$, $\tilde{p}_{ii} = 1 - \sum_{j \neq i} p_{ij}$. We still have $\pi_j \tilde{p}_{jk} = \pi_k \tilde{p}_{kj}$ and so (from Chapter 5), we know that the stationary probability of the chain $\tilde{P}$ is $\frac{1}{\pi(S)}\boldsymbol{\pi}$.

Further, the conductance of $\tilde{P}$ is at least ε by (8.3). Let $\mathbf{p}^{(t)}$ denote the probabilities of the chain $\tilde{P}$ at time t and let $\mathbf{a}$ denote the long term average, i.e.,

$$\mathbf{a} = \frac{1}{t_0} \sum_{t=0}^{t_0-1} \mathbf{p}^{(t)}.$$

From Theorem 5.12 of Chapter 5, $\left| \mathbf{a} - \frac{1}{\pi(S)}\boldsymbol{\pi} \right| \leq \frac{1}{100}$, Thus

$$a(S_1) \leq \frac{\pi(S_1)}{\pi(S)} + \frac{1}{100}. \quad (8.5)$$

Now, for each $i \in S$, $g_{it} \leq p_i^{(t)}$, since every run of the walk that never steps out of S has a corresponding run in $\tilde{P}$. So, $h_i \leq a_i$ and hence, $h(S_1) \leq a(S_1)$. Using (8.5) and (8.4):

$$h(S \setminus S_1) = h(S) - h(S_1) \geq h(S) - a(S_1) \geq 1 - O(\varepsilon) - \frac{\pi(S_1)}{\pi(S)} - \frac{1}{100} \geq \frac{\pi(S \setminus S_1)}{\pi(S)} - \frac{1}{50}.$$

But, $h(S \setminus S_1) \leq \frac{3}{4} \frac{\pi(S \setminus S_1)}{\pi(S)}$. Thus, $\pi(S \setminus S_1) \leq \frac{4}{50}$ proving the lemma. ■

Modularity clustering

Another way to partition a graph into communities is based on the concept of modularity. The method is popular for small graphs. Consider the partition of the vertices of a graph into communities. The modularity of the partition is defined to be the fraction of edges that lie within communities minus the expected number of edges that lie within communities in a random graph with the same degree distribution. Let A be the adjacency matrix of a graph, m the total number of edges, d_v the degree of vertex v , and let i index the communities defined by a partition of the vertices. Then, the modularity is given by

$$Q = \frac{1}{2m} \sum_i \sum_{v,w \in i} a_{vw} - \sum_i \sum_{v,w \in i} \frac{d_v}{2m} \frac{d_w}{2m}$$

Let e_{ij} denote the fraction of the total set of edges that connect communities i and j and let a_i denote the fraction of edges with ends in community i . Then

$$e_{ij} = \frac{1}{2m} \sum_{\substack{v \in i \\ w \in j}} a_{vw}$$

and

$$a_i = \frac{1}{2m} \sum_{v \in i} d_v$$

Write

$$\begin{aligned} Q &= \sum_i \sum_{v,w \in i} \left(\frac{1}{2m} a_{vw} - \frac{d_v}{2m} \frac{d_w}{2m} \right) \\ &= \sum_i (e_{ii} - a_i^2) \end{aligned}$$

The algorithm for finding communities works as follows. Start with each vertex in a community. Repeatedly merge the pair of communities that maximizes the change in Q . This can be done in time $O(m \log^2 n)$ where n is the number of vertices in the graph and m is the number of edges where $m \geq n$. The algorithm works well on small graphs but is inefficient for graphs with even a few thousand vertices.

Percolation clustering

Another clustering method that is useful in clustering nodes of a graph is called percolation clustering. Here one selects a value k and creates a new graph whose vertices correspond to k -cliques in the original graph. Edges in this new graph correspond to k -1 cliques connecting the k -cliques in the original graph. The connected components of this new graph are the clusters for the original graph.

8.11 Axioms for Clustering

Each clustering algorithm tries to optimize some criteria, like the sum of squared distances to the nearest cluster center, over all possible clustering's. We have seen many different optimization criteria in this chapter and many more are used. Now, we take a step back and ask what are the desirable properties of a clustering criterion and if there are criteria satisfying these properties. Our first result is negative. We present three seemingly desirable properties of a measure, and then show that no measure satisfies them. Next we argue that these requirements are too stringent and under more reasonable requirements, a slightly modified form of the sum of Euclidean distance squared between all pairs of points inside the same cluster is indeed a measure satisfying the desired properties.

8.11.1 An Impossibility Result

Let $A(d)$ denote the optimal clustering found by the clustering algorithm A using distance function d on a set S . The clusters of the clustering $A(d)$ form a partition Γ of S .

The first desirable property of a clustering algorithm is scale invariance. A clustering algorithm A is *scale invariant* if for any $\alpha > 0$, $A(d) = A(\alpha d)$. That is, multiplying all distances by some scale factor does not change the optimal clustering. In general, there could be ties for what the algorithm returns; in that case, we adopt the convention that $A(d) = A(\alpha d)$ really means for any clustering returned by A on distance d , it can also be returned by A on distance αd .

A clustering algorithm A is *rich* (full/complete) if for every partitioning Γ there exists a distance function d such that $A(d) = \Gamma$. That is, for any desired partitioning, we can find a set of distances so that the clustering algorithm returns the desired partitioning.

A clustering algorithm is *consistent* if increasing the distance between points in different clusters and reducing the distance between points in the same cluster does not change the clusters produced by the clustering algorithm.

If a clustering algorithm is consistent and $A(d) = \Gamma$, one can find a new distance function d' such that $A(d') = \Gamma$ where there are only two distances a and b . Here a is the distance between points within a cluster and b is the distance between points in different clusters. By consistency, we can reduce all distances within clusters and increase all distances between clusters there by getting two distances a and b with $a < b$ where a is the distance between points within a cluster and b is the distance between points in different clusters.

There exist natural clustering algorithms satisfying any two of the three axioms. The *single link clustering algorithm* starts with each point in a cluster by itself and then merges the two clusters that are closest. The process continues until some stopping condition is reached. One can view the process as the points being vertices of a graph and edges being

labeled by the distances between vertices. One merges the two vertices that are closest and merges parallel edges taking the distance of the merged edge to be the minimum of the two distances.

Theorem 8.15

1. *The single link clustering algorithm with the k -cluster stopping condition, stop when there are k clusters, satisfies scale-invariance and consistency. We do not get richness since we only get clustering's with k clusters.*
2. *The single link clustering algorithm with scale α stopping condition satisfies scale invariance and richness. The scale α stopping condition is to stop when the closest pair of clusters is of distance greater than or equal to $\alpha d_{\max}$ where $d_{\max}$ is the maximum pair wise distance. Here we do not get consistency. If we select one distance between clusters and increase it significantly until it becomes $d_{\max}$ and in addition $\alpha d_{\max}$ exceeds all other distances, the resulting clustering has just one cluster containing all of the points.*
3. *The single link clustering algorithm with the distance r stopping condition, stop when the inter-cluster distances are all at least r , satisfies richness and consistency; but not scale invariance.*

Proof: (1) Scale-invariance is easy to see. If one scales up all distances by a factor, then at each point in the algorithm, the same pair of clusters will be closest. The argument for consistency is more subtle. Since edges inside clusters of the optimal (final) clustering can only be decreased and since edges between clusters can only be increased, the edges that led to merges between any two clusters are less than any edge between the final clusters. Since the final number of clusters is fixed, these same edges will cause the same merges unless the merge has already occurred due to some other edge that was inside a final cluster having been shortened even more. No edge between two final clusters can cause a merge before all the above edges have been considered. At this time the final number of clusters has been reached and the process of merging has stopped.

(2) and (3) are straight forward. ■

Next, we show that no clustering algorithm can satisfy all three axioms. A distance function d , $(a-b)$ -conforms to a partition Γ if all points in a cluster are within distance a of each other and all points in distinct clusters are at least distance b apart. For a clustering algorithm A , the pair of numbers (a, b) is said to force the partition Γ if all distances functions d that $(a-b)$ conform to Γ have $A(d) = \Gamma$.

Associated with a clustering algorithm is a collection of allowable clustering's it can produce using different distance functions. We begin with a theorem stating that scale invariance and consistency imply that no allowable clustering can be a refinement of another allowable clustering.

Theorem 8.16 *If a clustering algorithm satisfies scale-invariance and consistency, then no two clustering's, one of which is a refinement of the other, can both be optimal clustering's returned by the algorithm.*

Proof: Suppose that the range of the clustering algorithm A contains two clustering's, Γ_0 and Γ_1 where Γ_0 is a refinement of Γ_1 . Modify the distance functions giving rise to Γ_0 and Γ_1 so that there are only two distinct distances a_0 and b_0 for Γ_0 and a_1 and b_1 for Γ_1 . Points within a cluster of Γ_0 are distance a_0 apart and points between clusters of Γ_0 are distance b_0 apart. The distances a_1 and b_1 play similar roles for Γ_1 .

Let a_2 be any number less than a_1 and choose ε such that $0 < \varepsilon < a_0 a_2 b_0^{-1}$. Let d be a new distance function where

$$d(i, j) = \begin{cases} \varepsilon & \text{if } i \text{ and } j \text{ are in the same cluster of } \Gamma_0 \\ a_2 & \text{if } i \text{ and } j \text{ are in different clusters of } \Gamma_0 \text{ but the same cluster of } \Gamma_1 \\ b_1 & \text{if } i \text{ and } j \text{ are in different clusters of } \Gamma_1 \end{cases}$$

From $a_0 < b_0$ it follows that $a_0 b_0^{-1} < 1$. Thus $\varepsilon < a_0 a_2 b_0^{-1} < a_2 < a_1$. Since both ε and a_2 are less than a_1 , it follows by consistency that $A(d) = \Gamma_1$. Let $\alpha = b_0 a_2^{-1}$. Since $\varepsilon < a_0 a_2 b_0^{-1}$ and $a_2 < a_1 < b_1$, which implies $a_2^{-1} > a_1^{-1} > b_1^{-1}$, it follows that

$$\alpha d(i, j) = \begin{cases} b_0 a_2^{-1} \varepsilon < b_0 a_2^{-1} a_0 a_2 b_0^{-1} = a_0 & \text{if } i \text{ and } j \text{ are in the same cluster of } \Gamma_0 \\ b_0 a_2^{-1} a_2 = b_0 & \text{if } i \text{ and } j \text{ are in different clusters of } \Gamma_0 \\ & \text{but the same cluster of } \Gamma_1 \\ b_0 a_2^{-1} b_1 > b_0 b_1^{-1} b_1 = b_0 & \text{if } i \text{ and } j \text{ are in different clusters of } \Gamma_1 \end{cases}$$

Thus, by consistency $A(\alpha d) = \Gamma_0$. But by scale invariance $A(\alpha d) = A(d) = \Gamma_1$, a contradiction. ■

Corollary 8.17 *For $n \geq 2$ there is no clustering function f that satisfies scale-invariance, richness, and consistency.*

It turns out that any collection of clustering's in which no clustering is a refinement of any other clustering in the collection is the range of a clustering algorithm satisfying scale invariance and consistency. To demonstrate this, we use the sum of pairs clustering algorithm. Given a collection of clustering's, the *sum of pairs clustering* algorithm finds the clustering that minimizes the sum of all distances between points in the same cluster over all clustering's in the collection.

Theorem 8.18 *Every collection of clustering's in which no clustering is the refinement of another is the range of a clustering algorithm A satisfying scale invariance and consistency.*

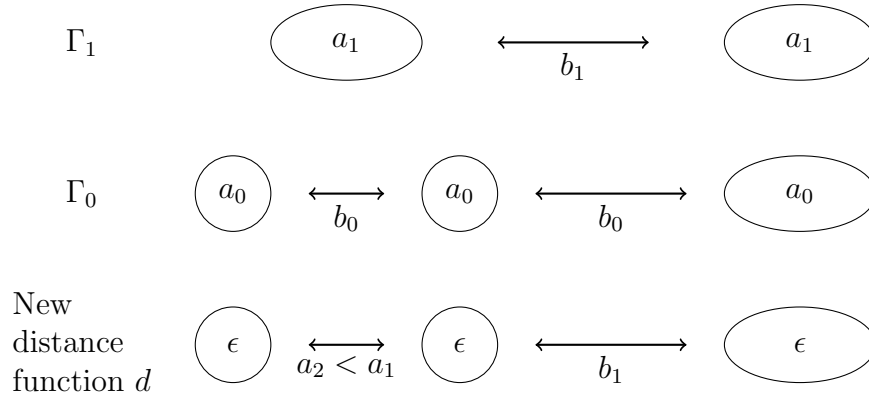


Figure 8.8: Illustration of the sets Γ_0, Γ_1 , and those for the distance function d .

Proof: We first show that the sum of pairs clustering algorithm satisfies scale invariance and consistency. Then we show that every collection of clustering's in which no cluster is a refinement of another can be achieved by a sum of pairs clustering algorithm.

Let A be the sum of pairs clustering algorithm. It is clear that A satisfies scale invariance since multiplying all distances by a constant, multiplies the total cost of each cluster by a constant and hence the minimum cost clustering is not changed.

To demonstrate that A satisfies consistency let d be a distance function and Γ the resulting clustering. Increasing the distance between pairs of points in different clusters of Γ does not affect the cost of Γ . If we reduce distances only between pairs of points in clusters of Γ then the cost of Γ is reduced as much or more than the cost of any other clustering. Hence Γ remains the lowest cost clustering.

Consider a collection of clustering's in which no cluster is a refinement of another. It remains to show that every clustering in the collection is in the range of A . In sum of pairs clustering, the minimum is over all clustering's in the collection. We now show for any clustering Γ how to assign distances between pairs of points so that A returns the desired clustering. For pairs of points in the same cluster assign a distance of $1/n^3$. For pairs of points in different clusters assign distance one. The cost of the clustering Γ is less than one. Any clustering that is not a refinement of Γ has cost at least one. Since there are no refinements of Γ in the collection it follows that Γ is the minimum cost clustering. ■

Note that one may question both the consistency axiom and the richness axiom. The following are two possible objections to the consistency axiom. Consider the two clusters in Figure 8.9. If one reduces the distance between points in cluster B , they might get an arrangement that should be three clusters instead of two.



Figure 8.9: Illustration of the objection to the consistency axiom. Reducing distances between points in a cluster may suggest that the cluster be split into two.

The other objection, which applies to both the consistency and the richness axioms, is that they force many unrealizable distances to exist. For example, suppose the points were in Euclidean d space and distances were Euclidean. Then, there are only nd degrees of freedom. But the abstract distances used here have $O(n^2)$ degrees of freedom since the distances between the $O(n^2)$ pairs of points can be specified arbitrarily. Unless d is about n , the abstract distances are too general. The objection to richness is similar. If for n points in Euclidean d space, the clusters are formed by hyper planes each cluster may be a Voronoi cell or some other polytope, then as we saw in the theory of VC dimensions Section ?? there are only $\binom{n}{d}$ interesting hyper planes each defined by d of the n points. If k clusters are defined by bisecting hyper planes of pairs of points, there are only n^{dk^2} possible clustering's rather than the 2^n demanded by richness. If d and k are significantly less than n , then richness is not reasonable to demand. In the next section, we will see a possibility result to contrast with this impossibility theorem.

The k -means clustering algorithm is one of the most widely used clustering algorithms. We now show that any centroid based algorithm such as k -means does not satisfy the consistency axiom.

Theorem 8.19 *A centroid based clustering such as k -means does not satisfy the consistency axiom.*

Proof: The cost of a cluster is $\sum_i (\mathbf{x}_i - \mathbf{u})^2$, where u is the centroid. An alternative way to compute the cost of the cluster if the distances between pairs of points in the cluster are known is to compute $\frac{1}{n} \sum_{i \neq j} (\mathbf{x}_i - \mathbf{x}_j)^2$ where n is the number of points in the cluster. For a proof see Lemma 8.2. Consider seven points, a point $\mathbf{y}$ and two sets of three points each, called X_0 and X_1 . Let the distance from $\mathbf{y}$ to each point in $X_0 \cup X_1$ be $\sqrt{5}$ and let all other distances between pairs of points be $\sqrt{2}$. These distances are achieved by placing each point of X_0 and X_1 a distance one from the origin along a unique coordinate and placing $\mathbf{y}$ at distance two from the origin along another coordinate. Consider a clustering with two clusters (see Figure 8.9). The cost depends only on how many points are grouped with $\mathbf{y}$. Let that number be m . The cost is

$$\frac{1}{m+1} \left[2 \binom{m}{2} + 5m \right] + \frac{2}{6-m} \binom{6-m}{2} = \frac{8m+5}{m+1}$$

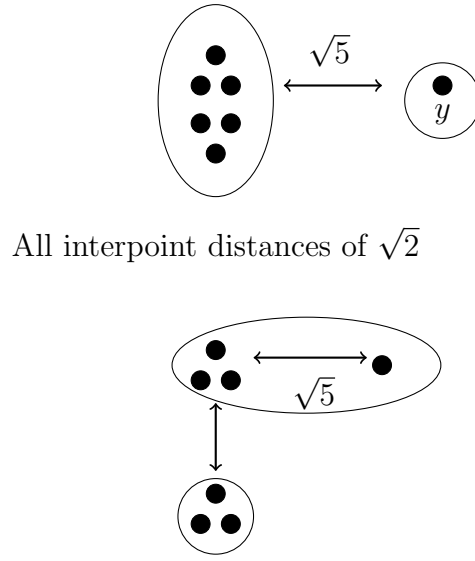


Figure 8.10: Example illustrating k -means does not satisfy the consistency axiom.

which has its minimum at $m = 0$. That is, the point $\mathbf{y}$ is in a cluster by itself and all other points are in a second cluster.

If we now shrink the distances between points in X_0 and points in X_1 to zero, the optimal clustering changes. If the clusters were $X_0 \cup X_1$ and $\mathbf{y}$, then the distance would be $9 \times 2 = 18$ whereas if the clusters are $X_0 \cup \{\mathbf{y}\}$ and X_1 , the distance would be only $3 \times 5 = 15$. Thus, the optimal clustering is $X_0 \cup \{\mathbf{y}\}$ and X_1 . Hence k -means does not satisfy the consistency axiom since shrinking distances within clusters changes the optimal clustering. ■

5 Relaxing the axioms

Given that no clustering algorithm can satisfy scale invariance, richness, and consistency, one might want to relax the axioms in some way. Then one gets the following results.

1. Single linkage with a distance stopping condition satisfies a relaxed scale-invariance property that states that for $\alpha > 1$, then $f(\alpha d)$ is a refinement of $f(d)$.

2. Define *refinement consistency* to be that shrinking distances within a cluster or expanding distances between clusters gives a refinement of the clustering. Single linkage with α stopping condition satisfies scale invariance, refinement consistency and richness except for the trivial clustering of all singletons.

8.11.2 A Satisfiable Set of Axioms

In this section, we propose a different set of axioms that are reasonable for distances between points in Euclidean space and show that the clustering measure, the sum of squared distances between all pairs of points in the same cluster, slightly modified, is consistent with the new axioms. We assume through the section that points are in Euclidean d -space. Our three new axioms follow.

We say that a clustering algorithm satisfies the *consistency condition* if, for the clustering produced by the algorithm on a set of points, moving a point so that its distance to any point in its own cluster is not increased and its distance to any point in a different cluster is not decreased, then the algorithm returns the same clustering after the move.

Remark: Although it is not needed in the sequel, it is easy to see that for an infinitesimal perturbation dx of x , the perturbation is consistent if and only if each point in the cluster containing x lies in the half space through x with dx as the normal and each point in a different cluster lies in the other half space.

An algorithm is *scale-invariant* if multiplying all distances by a positive constant does not change the clustering returned.

An algorithm has the *richness* property if for any set K of k distinct points in the ambient space, there is some placement of a set S of n points to be clustered so that the algorithm returns a clustering with the points in K as centers. So there are k clusters, each cluster consisting of all points of S closest to one particular point of K .

We will show that the following algorithm satisfies these three axioms.

Balanced k -means algorithm

Among all partitions of the input set of n points into k sets, each of size n/k , return the one that minimizes the sum of squared distances between all pairs of points in the same cluster.

Theorem 8.20 *The balanced k -means algorithm satisfies the consistency condition, scale invariance, and the richness property.*

Proof: Scale invariance is obvious. Richness is also easy to see. Just place n/k points of S to coincide with each point of K . To prove consistency, define the *cost* of a cluster T to be the sum of squared distances of all pairs of points in T .

Suppose $S_1, S_2, \dots, S_k$ is an optimal clustering of S according to the balanced k -means algorithm. Move a point $x \in S_1$ to z so that its distance to each point in S_1 is non increasing and its distance to each point in $S_2, S_3, \dots, S_k$ is non decreasing. Suppose

$T_1, T_2, \dots, T_k$ is an optimal clustering after the move. Without loss of generality assume $z \in T_1$. Define $\tilde{T}_1 = (T_1 \setminus \{z\}) \cup \{x\}$ and $\tilde{S}_1 = (S_1 \setminus \{x\}) \cup \{z\}$. Note that $\tilde{T}_1, T_2, \dots, T_k$ is a clustering before the move, although not necessarily an optimal clustering. Thus

$$\text{cost}(\tilde{T}_1) + \text{cost}(T_2) + \dots + \text{cost}(T_k) \geq \text{cost}(S_1) + \text{cost}(S_2) + \dots + \text{cost}(S_k).$$

If $\text{cost}(T_1) - \text{cost}(\tilde{T}_1) \geq \text{cost}(\tilde{S}_1) - \text{cost}(S_1)$ then

$$\text{cost}(T_1) + \text{cost}(T_2) + \dots + \text{cost}(T_k) \geq \text{cost}(\tilde{S}_1) + \text{cost}(S_2) + \dots + \text{cost}(S_k).$$

Since $T_1, T_2, \dots, T_k$ is an optimal clustering after the move, so also must be $\tilde{S}_1, S_2, \dots, S_k$ proving the theorem.

It remains to show that $\text{cost}(T_1) - \text{cost}(\tilde{T}_1) \geq \text{cost}(\tilde{S}_1) - \text{cost}(S_1)$. Let u and v stand for elements other than x and z in S_1 and T_1 . The terms $|u - v|^2$ are common to T_1 and $\tilde{T}_1$ on the left hand side and cancel out. So too on the right hand side. So we need only prove

$$\sum_{u \in T_1} (|z - u|^2 - |x - u|^2) \geq \sum_{u \in S_1} (|z - u|^2 - |x - u|^2).$$

For $u \in S_1 \cap T_1$, the terms appear on both sides, and we may cancel them, so we are left to prove

$$\sum_{u \in T_1 \setminus S_1} (|z - u|^2 - |x - u|^2) \geq \sum_{u \in S_1 \setminus T_1} (|z - u|^2 - |x - u|^2)$$

which is true because by the movement of x to z , each term on the left hand side is non negative and each term on the right hand side is non positive. ■

8.12 Exercises

Exercise 8.1 Construct examples where using distances instead of distance squared gives bad results for Gaussian densities. For example, pick samples from two 1-dimensional unit variance Gaussians, with their centers 10 units apart. Cluster these samples by trial and error into two clusters, first according to k -means and then according to the k -median criteria. The k -means clustering should essentially yield the centers of the Gaussians as cluster centers. What cluster centers do you get when you use the k -median criterion?

Exercise 8.2 Let $v = (1, 3)$. What is the L_1 norm of v ? The L_2 norm? The square of the L_1 norm?

Exercise 8.3 Show that in 1-dimension, the center of a cluster that minimizes the sum of distances of data points to the center is in general not unique. Suppose we now require the center also to be a data point; then show that it is the median element (not the mean). Further in 1-dimension, show that if the center minimizes the sum of squared distances to the data points, then it is unique.

Exercise 8.4 Construct a block diagonal matrix A with three blocks of size 50. Each matrix element in a block has value $p = 0.7$ and each matrix element not in a block has value $q = 0.3$. generate a 150×150 matrix B of random numbers in the range $[0, 1]$. If $b_{ij} \geq a_{ij}$ replace a_{ij} with the value one. Otherwise replace a_{ij} with value zero. The rows of A have three natural clusters. Permute the rows and columns of A so the first 50 rows do not form the first cluster, the next 50 the second cluster, and the last 50 the third cluster.

1. Apply the k -mean algorithm to A with $k = 3$. Do you find the correct clusters?
2. Apply the k -means algorithm to A for $1 \leq k \leq 10$. Plot the value of the sum of squares to the cluster centers versus k . Was three the correct value for k ?

Exercise 8.5 Let M be a $k \times k$ matrix whose elements are numbers in the range $[0, 1]$. A matrix entry close to one indicates that the row and column of the entry correspond to closely related items and an entry close to zero indicates unrelated entities. Develop an algorithm to match each row with a closely related column where a column can be matched with only one row.

Exercise 8.6 The simple greedy algorithm of Section 8.3 assumes that we know the clustering radius r . Suppose we do not. Describe how we might arrive at the correct r ?

Exercise 8.7 For the k -median problem, show that there is at most a factor of two ratio between the optimal value when we either require all cluster centers to be data points or allow arbitrary points to be centers.

Exercise 8.8 For the k -means problem, show that there is at most a factor of four ratio between the optimal value when we either require all cluster centers to be data points or allow arbitrary points to be centers.

Exercise 8.9 Consider clustering points in the plane according to the k -median criterion, where cluster centers are required to be data points. Enumerate all possible clustering's and select the one with the minimum cost. The number of possible ways of labeling n points, each with a label from $\{1, 2, \dots, k\}$ is k^n which is prohibitive. Show that we can find the optimal clustering in time at most a constant times $\binom{n}{k} + k^2$. Note that $\binom{n}{k} \leq n^k$ which is much smaller than k^n when $k \ll n$.

Exercise 8.10 Suppose in the previous exercise, we allow any point in space (not necessarily data points) to be cluster centers. Show that the optimal clustering may be found in time at most a constant times n^{2k^2} .

Exercise 8.11 Corollary 8.3 shows that for a set of points $\{a_1, a_2, \dots, a_n\}$, there is a unique point x , namely their centroid, which minimizes $\sum_{i=1}^n |a_i - x|^2$. Show examples where the x minimizing $\sum_{i=1}^n |a_i - x|$ is not unique. (Consider just points on the real line.) Show examples where the x defined as above are far apart from each other.

Exercise 8.12 Let $\{\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_n\}$ be a set of unit vectors in a cluster. Let $\mathbf{c} = \frac{1}{n} \sum_{i=1}^n \mathbf{a}_i$ be the cluster centroid. The centroid $\mathbf{c}$ is not in general a unit vector. Define the similarity between two points $\mathbf{a}_i$ and $\mathbf{a}_j$ as their dot product. Show that the average cluster similarity $\frac{1}{n^2} \sum_{i,j} \mathbf{a}_i \mathbf{a}_j^T$ is the same whether it is computed by averaging all pairs or computing the average similarity of each point with the centroid of the cluster.

Exercise 8.13 For some synthetic data estimate the number of local minima for k -means by using the birthday estimate. Is your estimate an unbiased estimate of the number? an upper bound? a lower bound? Why?

Exercise 8.14 Examine the example in Figure and discuss how to fix it. Optimizing according to the k -center or k -median criteria would seem to produce clustering B while clustering A seems more desirable.

Exercise 8.15 Prove that for any two vectors $\mathbf{a}$ and $\mathbf{b}$, $|\mathbf{a} - \mathbf{b}|^2 \geq \frac{1}{2}|\mathbf{a}|^2 - |\mathbf{b}|^2$.

Exercise 8.16 Let A be an $n \times d$ data matrix, B its best rank k approximation, and C the optimal centers for k -means clustering of rows of A . How is it possible that $\|A - B\|_F^2 < \|A - C\|_F^2$?

Exercise 8.17 Suppose S is a finite set of points in space with centroid $\mu(S)$. If a set T of points is added to S , show that the centroid $\mu(S \cup T)$ of $S \cup T$ is at distance at most $\frac{|T|}{|S|+|T|} |\mu(T) - \mu(S)|$ from μ .

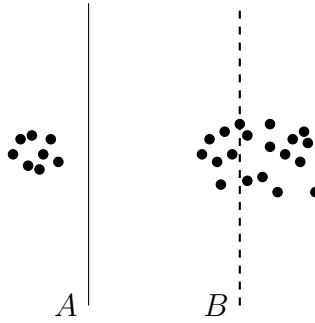


Figure 8.11: insert caption

Exercise 8.18 What happens if we relax this restriction, for example, if we allow for S , the entire set?

Exercise 8.19 Given the graph $G = (V, E)$ of a social network where vertices represent individuals and edges represent relationships of some kind, one would like to define the concept of a community. A number of deferent definitions are possible.

1. A subgraph $S = (V_S, E_S)$ whose density $\frac{E_S}{V_S^2}$ is greater than that of the graph $\frac{E}{V^2}$.
2. A subgraph S with a low conductance like property such as the number of graph edges leaving the subgraph normalized by the minimum size of S or $V - S$ where size is measured by the sum of degrees of vertices in S or in $V - S$.
3. A subgraph that has more internal edges than in a random graph with the same degree distribution.

Which would you use and why?

Exercise 8.20 A stochastic matrix is a matrix with non negative entries in which each row sums to one. Show that for a stochastic matrix, the largest eigenvalue is one. Show that the eigenvalue has multiplicity one if and only if the corresponding Markov Chain is connected.

Exercise 8.21 Show that if P is a stochastic matrix and π satisfies $\pi_i p_{ij} = \pi_j p_{ji}$, then for any left eigenvector $\mathbf{v}$ of P , the vector $\mathbf{u}$ with components $u_i = \frac{v_i}{\pi_i}$ is a right eigenvector with the same eigenvalue.

Exercise 8.22 In Theorem (??), how can one clustering $C^{(0)}$ be close to any proper clustering? What if there are several proper clusterings?

Exercise 8.23 Give an example of a clustering problem where the clusters are not linearly separable in the original space, but are separable in a higher dimensional space.

Hint: Look at the example for Gaussian kernels in the chapter on learning.

Exercise 8.24 The Gaussian kernel maps points to a higher dimensional space. What is this mapping?

Exercise 8.25 Agglomerative clustering requires that one calculate the distances between all pairs of points. If the number of points is a million or more, then this is impractical. One might try speeding up the agglomerative clustering algorithm by maintaining a 100 clusters at each unit of time. Start by randomly selecting a hundred points and place each point in a cluster by itself. Each time a pair of clusters is merged randomly select one of the remaining data points and create a new cluster containing that point. Suggest some other alternatives.

Exercise 8.26 Let A be the adjacency matrix of an undirected graph. Let $d(S, S) = \frac{A(S, S)}{|S|}$ be the density of the subgraph induced by the set of vertices S . Prove that $d(S, S)$ is the average degree of a vertex in S .

Exercise 8.27 Suppose A is a matrix with non negative entries. Show that $A(S, T)/(|S||T|)$ is maximized by the single edge with highest a_{ij} .

Exercise 8.28 Suppose A is a matrix with non negative entries and

$$\sigma_1(A) = \mathbf{x}^T A \mathbf{y} = \sum_{i,j} x_i a_{ij} y_j, \quad |\mathbf{x}| = |\mathbf{y}| = 1.$$

Zero out all x_i less than $1/2\sqrt{n}$ and all y_j less than $1/2\sqrt{d}$. Show that the loss is no more than $1/4^{\text{th}}$ of $\sigma_1(A)$.

Exercise 8.29 Consider other measures of density such as $\frac{A(S, T)}{|S|^\rho |T|^\rho}$ for different values of ρ . Discuss the significance of the densest subgraph according to these measures.

Exercise 8.30 Let A be the adjacency matrix of an undirected graph. Let M be the matrix whose ij^{th} element is $a_{ij} - \frac{d_i d_j}{2m}$. Partition the vertices into two groups S and $\bar{S}$. Let s be the indicator vector for the set S and let $\bar{s}$ be the indicator variable for $\bar{S}$. Then $s^T M s$ is the number of edges in S above the expected number given the degree distribution and $s^T M \bar{s}$ is the number of edges from S to $\bar{S}$ above the expected number given the degree distribution. Prove that if $s^T M s$ is positive $s^T M \bar{s}$ must be negative.

Exercise 8.31 Which of the three axioms, scale invariance, richness, and consistency are satisfied by the following clustering algorithms.

1. k -means
2. Spectral Clustering.

Exercise 8.32 (Research Problem): What are good measures of density that are also effectively computable? Is there empirical/theoretical evidence that some are better than others?

9 Topic Models, Hidden Markov Process, Graphical Models, and Belief Propagation

In the chapter on learning and VC dimension, we saw many model-fitting problems. There, we were given labeled data and simple classes of functions - half-spaces, support vector machines, etc. The problem was to fit the best model from a class of functions to the data. Model fitting is of course more general and in this chapter, we discuss some useful models. These general models are often computationally infeasible, in the sense that they do not admit provably efficient algorithms. Nevertheless, data often falls into special cases of these models that can be solved efficiently.

9.1 Topic Models

A *topic model* is a model for representing a large collection of documents. In the abstract, each document is viewed as a combination of topics and each topic has a set of word frequencies. For a collection of news articles over a period, the topics may be politics, sports, science, etc. For the topic politics, the words like “president”, “election” may have high frequencies and for the topic sports, words like “batter” and “goal” may have high frequencies. A news item document may be 60% on politics and 40% on sports. The word frequencies in the document will be convex combinations of word frequencies for the topics, politics and sports, with weights 0.6 and 0.4 respectively. We describe this more formally with vectors and matrices.

Each document is viewed as a “bag of words”. We disregard the order and context in which each word occurs in the document and instead only list the frequency of occurrences of each term. Frequency is the number of occurrences of the term divided by the total number of all terms in the document. Discarding context information may seem wasteful, but this approach works well in practice and is widely used. Each document is an n -dimensional vector where n is the total number of different terms in all the documents in the collection. Each component of the vector is the frequency of a particular term in the document. Terms are words or phrases. Not all words are chosen as terms; articles, simple verbs, and pronouns like “a”, “is”, and “it” may be ignored. Represent the collection of documents by a $n \times m$ matrix A , called the *term-document* matrix, with one column per document in the collection. The topic model hypothesizes that there are r topics and each of the m documents is a combination of topics. The number of topics r is usually much smaller than the number of terms n . So corresponding to each document, there is a vector with r components telling us the fraction of the document that is on each of the topics. In the example above, this vector will have 0.6 in the component for politics and 0.4 in the component for sports. Arrange these vectors as the columns of a $r \times m$ matrix B , called the *topic-document* matrix. There is a third matrix C which is $n \times r$. Each column of C corresponds to a topic; each component of the column gives the frequency of a term in that topic. In the simplest model, the term frequencies in documents are exactly combinations of term frequencies in the various topics that make up the document. So,

a_{ij} , the frequency of the i^{th} term in the j^{th} document is the sum over all topics l of the fraction of document j which is on topic l times the frequency of term i in topic l . In matrix notation,

$$A = CB.$$

Pictorially, we can represent this as:

$$\begin{array}{c} \text{T} \\ \text{E} \\ \text{R} \\ \text{M} \end{array} \left(\begin{array}{c} \text{DOCUMENT} \\ \\ \\ \end{array} \right) = \begin{array}{c} \text{T} \\ \text{E} \\ \text{R} \\ \text{M} \end{array} \left(\begin{array}{c} \text{TOPIC} \\ \\ \\ \end{array} \right) \begin{array}{c} \text{T} \\ \text{O} \\ \text{P} \\ \text{I} \\ \text{C} \end{array} \left(\begin{array}{c} \text{DOCUMENT} \\ \\ \\ \end{array} \right)$$

This model is too simple to be realistic since the frequency of each term in the document is unlikely to be exactly what is given by the equation $A = CB$. So, a more sophisticated stochastic model is used in practice.

From the document collection we observe the $n \times m$ matrix A . Can we find C and B such that $A = CB$? The top r singular vectors from a singular value decomposition of A give a factorization CB . But there are additional constraints stemming from the fact that frequencies of terms in one particular topic are nonnegative reals summing to one and from the fact that the fraction of each topic a particular document is on are also nonnegative reals summing to one. Altogether the constraints are:

1. $A = CB$.
2. The entries of B and C are all non-negative.
3. Each column of B and each column of C sums to one.

Given the first two conditions we can achieve the third by multiplying the i^{th} column of C by a positive real number and dividing the i^{th} row of B by the same real number without violating $A = CB$. By doing this, one may assume that each column of C sums to one. Now $a_{ij} = \sum_l c_{il}b_{lj}$ implies $\sum_i a_{ij} = \sum_{i,l} c_{il}b_{lj} = \sum_l b_{lj}$. But, $\sum_i a_{ij}$ is the total frequency of all terms in document j and must equal one. Thus, the columns of B also sum to one.

The problem can be posed as one of factoring the given matrix A into the product of two matrices with nonnegative entries called nonnegative matrix factorization.

Nonnegative matrix factorization - NMF Given an $n \times m$ matrix A and an integer r , determine whether there is a factorization of A into XY where, X is an $n \times r$

matrix with nonnegative entries and Y is $r \times m$ matrix with nonnegative entries and if so, find such a factorization.

Nonnegative matrix factorization is a more general problem than topic modeling and there are many heuristic algorithms to solve the problem. But in general, they suffer from one of two problems, they can get stuck at local optima that are not solutions or take exponential time. In fact, the general NMF problem is NP-hard. In practice, often r is much smaller than n and m . We show first that while the NMF problem as formulated above is a nonlinear problem in $r(n + m)$ unknown entries of X and Y , it can be reformulated as a nonlinear problem with just $2r^2$ unknowns under the simple nondegeneracy assumption that A has rank r . Think of r as say, 25, while n and m are in the tens of thousands to see why this is useful.

Lemma 9.1 *If A has rank r , then the NMF problem can be formulated as a problem with $2r^2$ unknowns.*

Proof: If $A = XY$, then each row of A is a linear combination of the rows of Y . So we have that the space spanned by the rows of A must be contained in the space spanned by the rows of Y . The latter space has dimension at most r , while the former has dimension r . So they must be equal. Thus, every row of Y must be a linear combination of the rows of A . Choose any set of r independent rows of A to form a $r \times m$ matrix A_1 . Then $Y = SA_1$ for some $r \times r$ matrix S . By exactly analogous reasoning, if A_2 is a $n \times r$ matrix of r independent columns of A , there is a $r \times r$ matrix T such that $X = A_2T$. Now we can easily cast NMF in terms of the unknowns S and T .

$$A = A_2TSA_1 \quad (SA_1)_{ij} \geq 0 \quad (A_2T)_{kl} \geq 0 \quad \forall i, j, k, l.$$

■

It remains to solve the nonlinear problem in the $2r^2$ variables. There is a classical algorithm that solves such problems in time exponential only in r^2 and polynomial in the other parameters. In fact, there is a logical theory, called the theory of reals of which this is a special case and any problem in this theory can be solved in time exponential only in the number of variables. We do not give details here.

Besides the special case when r is small, there is another important case of NMF in the topic modeling application that can be solved. This is the case when there are *anchor terms*. An anchor term for a topic is a term that occurs in the topic and does not occur in any other topic. For example, the term “batter” may an anchor term for the topic baseball and “election” for the topic politics. Consider the case when each topic has an anchor term. In matrix notation, this assumes that for each column of the term-topic matrix C , there is a row whose sole nonzero entry is in that column. In this case, it is easy to see that each row of the topic-document matrix B has a scaled copy of it occurring

as a row of the given term-document matrix A . Here is an illustrative diagram:

$$\begin{pmatrix} 0.3 \times b_4 \\ \\ A \\ 0.2 \times b_2 \end{pmatrix} = \begin{matrix} \text{election} \\ \\ \\ \text{batter} \end{matrix} \begin{pmatrix} 0 & 0 & 0 & 0.3 \\ & & C & \\ 0 & 0.2 & 0 & 0 \end{pmatrix} \begin{pmatrix} b_1 \\ b_2 \\ b_3 \\ \vdots \\ b_n \end{pmatrix}.$$

If we knew which rows of A are copies of rows of B , call these special rows of A , we would be done since then we can find B and once B is known, we can solve linear equations and inequalities ($A = CB$; $c_{ij} \geq 0$) to get C . The following lemma shows that after making one modification, we can find the rows of A that are special. Suppose a row of B is a nonnegative linear combination of the other rows of B . Then, we can eliminate that row of B as well as the corresponding column of C , suitably modifying the other columns of C , and still maintain $A = CB$. For example, if row 5 of B equals 4 times row 3 of B plus 3 times row 6 of B , then we can delete row 5 of B , then add 4 times column 5 of C to column 3 of C and add 3 times column 5 of C to column 6 of C and delete column 5 of C .

After repeating this, we may assume that each row of B is positively independent of the other rows of B , i.e., it cannot be expressed as a nonnegative linear combination of the other rows. We still have a scaled copy of each row of B in A . Further, the other rows of A are all nonnegative linear combinations of rows of B and thus are nonnegative linear combinations of the special rows of A .

Lemma 9.2 *Suppose A has a factorization $A = CB$, where the rows of B are positively independent and for each column of C , there is a row that has its sole nonzero entry in that column. Then there is a scaled copy of each row of B in A and furthermore, the rows of A that are scaled copies of rows of B are precisely the rows of A that are positively independent of other rows of A . These rows can be identified by solving a linear program, one program per row.*

Proof: What remains to prove is that by solving n linear programming problems, we can identify the set of special rows of A . For each row of A , we need to check if it is positively independent of all other rows. Denote by $\mathbf{a}_i$ the i th row of A . Then, the i th row is positively dependent upon the others if and only if there are real numbers $x_1, x_2, \dots, x_{i-1}, x_{i+1}, \dots, x_n$ such that

$$\sum_{j \neq i} x_j \mathbf{a}_j = \mathbf{a}_i \quad x_j \geq 0.$$

This is a linear program. ■

As we remarked earlier, the equation $A = CB$ will not hold exactly. A more practical model views A as a matrix of probabilities rather than exact frequencies. In this model, each document is generated by picking its terms in independent trials. Each trial for document j picks Term 1 with probability a_{1j} ; Term 2 with probability a_{2j} , etc. We are not given entire documents; instead we are given s independent trials for each document. Our job is still to find C and B . We will not discuss here the details of either the model or the algorithms. But in this new situation, if we assume the existence of anchor terms, algorithms are known to find C and B even with fairly small number s of trials.

At the heart of such an algorithm is the following problem:

Approximate NMF Given a $n \times m$ matrix A and the promise that there is a $n \times r$ matrix C and a $r \times m$ matrix B , both with nonnegative entries, such that $\|A - CB\|_F \leq \Delta$, find C' and B' of the same dimensions, with nonnegative entries such that $\|A - C'B'\|_F \leq \Delta'$.

Here, Δ' is related to Δ and if the promise does not hold, the algorithm is allowed to return any answer.

Now for the case when anchor words exist, this reduces to the problem of finding which rows of A have the property that no point close to the row is positively dependent on other rows. It is easy to write the statement that there is a vector $\mathbf{y}$ close to $\mathbf{a}_i$ which is positively dependent on the other rows as a convex program:

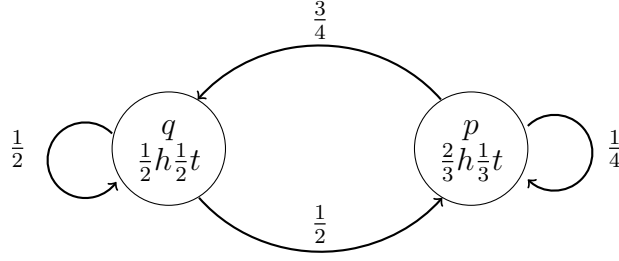
$$\exists x_1, x_2, \dots, x_{i-1}, x_{i+1}, \dots, x_n \text{ such that } \left| \sum_{j \neq i} x_j \mathbf{a}_j - \mathbf{a}_i \right| \leq \varepsilon.$$

$|\sum_{j \neq i} x_j \mathbf{a}_j - \mathbf{a}_i|$ is convex function of x_j and hence this problem can be solved efficiently.

9.2 Hidden Markov Model

A *hidden Markov model*, HMM, consists of a finite state graph with a transition between each pair of states. There is an initial probability distribution α on the states and a transition probability a_{ij} associated with the transition from state i to state j where $\sum_{j=1}^n a_{ij} = 1$. Each state has a probability distribution $p(O, i)$ giving the probability of outputting the symbol O in state i . A transition consists of two components. A state transition to a new state followed by the output of a symbol. The HMM starts by selecting a start state according to the distribution α and outputting a symbol.

Example: An example of a HMM is the graph with two states q and p illustrated below.



The initial distribution is $\alpha(q) = 1$ and $\alpha(p) = 0$. At each step a change of state occurs followed by the output of heads or tails with probability determined by the new state. ■

We consider three problems in increasing order of difficulty. First, given a HMM what is the probability of a given output sequence? Second, given a HMM and an output sequence, what is the most likely sequence of states? And third, knowing that the HMM has at most n states and given an output sequence, what is the most likely HMM?

How probable is an output sequence

Given a HMM, how probable is the output sequence $O = O_0O_1O_2 \cdots O_T$ of length $T+1$? To determine this, calculate for each state i and each initial segment of the sequence of observations, $O_0O_1O_2 \cdots O_t$ of length $t+1$, the probability of observing $O_0O_1O_2 \cdots O_t$ ending in state i . This is done by a dynamic programming algorithm starting with $t = 0$ and increasing t . For $t = 0$ there have been no transitions. Thus, the probability of observing O_0 ending in state i is the initial probability of starting in state i times the probability of observing O_0 in state i . The probability of observing $O_0O_1O_2 \cdots O_t$ ending in state i is the sum of the probabilities over all states j of observing $O_0O_1O_2 \cdots O_{t-1}$ ending in state j times the probability of going from state j to state i and observing O_t . The time to compute the probability of a sequence of length T when there are n states is $O(n^2T)$. The factor n^2 comes from the calculation for each time unit of the contribution from each possible previous state to the probability of each possible current state. The space complexity is $O(n)$ since one only needs to remember the probability of reaching each state for the most recent value of t .

Algorithm to calculate the probability of the output sequence

The probability, $\text{Prob}(O_0O_1 \cdots O_T, i)$ of the output sequence $O_0O_1 \cdots O_T$ ending in state i is given by

$$\text{Prob}(O_0, i) = \alpha(i)p(O_0, i)$$

for $t = 1$ to T

$$\text{Prob}(O_0O_1 \cdots O_t, i) = \sum_j \text{Prob}(O_0O_1 \cdots O_{t-1}, j)a_{ij}p(O_{t+1}, i)$$

Example: What is the probability of the sequence hhht by the HMM in the above example?

$t = 3$	$\frac{3}{32} \frac{1}{2} \frac{1}{2} + \frac{5}{72} \frac{3}{4} \frac{1}{2} = \frac{19}{384}$	$\frac{3}{32} \frac{1}{2} \frac{1}{3} + \frac{5}{72} \frac{1}{4} \frac{1}{3} = \frac{37}{64 \times 27}$
$t = 2$	$\frac{1}{8} \frac{1}{2} \frac{1}{2} + \frac{1}{6} \frac{3}{4} \frac{1}{2} = \frac{3}{32}$	$\frac{1}{8} \frac{1}{2} \frac{2}{3} + \frac{1}{6} \frac{1}{4} \frac{2}{3} = \frac{5}{72}$
$t = 1$	$\frac{1}{2} \frac{1}{2} \frac{1}{2} = \frac{1}{8}$	$\frac{1}{2} \frac{1}{2} \frac{2}{3} = \frac{1}{6}$
$t = 0$	$\frac{1}{2}$	0
	q	p

For $t = 0$, the q entry is $1/2$ since the probability of being in state q is one and the probability of outputting heads is $\frac{1}{2}$. The entry for p is zero since the probability of starting in state p is zero. For $t = 1$, the q entry is $\frac{1}{8}$ since for $t = 0$ the q entry is $\frac{1}{2}$ and in state q the HMM goes to state q with probability $\frac{1}{2}$ and outputs heads with probability $\frac{1}{2}$. The p entry is $\frac{1}{6}$ since for $t = 0$ the q entry is $\frac{1}{2}$ and in state q the HMM goes to state p with probability $\frac{1}{2}$ and outputs heads with probability $\frac{2}{3}$. For $t = 2$, the q entry is $\frac{3}{32}$ which consists of two terms. The first term is the probability of ending in state q at $t = 1$ times the probability of staying in q and outputting h . The second is the probability of ending in state p at $t = 1$ times the probability of going from state p to state q and outputting h .

From the table, the probability of producing the sequence hhht is $\frac{19}{384} + \frac{37}{1728} = 0.0709$. ■

The most likely sequence of states - the Viterbi algorithm

Given a HMM and an observation $O = O_0 O_1 \dots O_T$, what is the most likely sequence of states? The solution is given by the Viterbi algorithm, which is a slight modification to the dynamic programming algorithm just given for determining the probability of an output sequence. For $t = 0, 1, 2, \dots, T$ and for each state i , calculate the probability of the most likely sequence of states to produce the output $O_0 O_1 O_2 \dots O_t$ ending in state i . For each value of t , calculate the most likely sequence of states by selecting over all states j the most likely sequence producing $O_0 O_1 O_2 \dots O_t$ and ending in state i consisting of the most likely sequence producing $O_0 O_1 O_2 \dots O_{t-1}$ ending in state j followed by the transition from j to i producing O_t . Note that in the previous example, we added the probabilities of each possibility together. Now we take the maximum and also record where the maximum came from. The time complexity is $O(n^2 T)$ and the space complexity is $O(nT)$. The space complexity bound is argued as follows. In calculating the probability of the most likely sequence of states that produces $O_0 O_1 \dots O_t$ ending in state i , we remember the previous state j by putting an arrow with edge label t from i to j . At the end, can find the most likely sequence by tracing backwards as is standard for dynamic programming algorithms.

Example: For the earlier example what is the most likely sequence of states to produce the output hhht?

$t = 3$	$\max\{\frac{1}{48}\frac{1}{2}\frac{1}{2}, \frac{1}{24}\frac{3}{4}\frac{1}{2}\} = \frac{1}{64} \quad q \text{ or } p$	$\max\{\frac{3}{48}\frac{1}{2}\frac{1}{3}, \frac{1}{24}\frac{1}{4}\frac{1}{3}\} = \frac{1}{96} \quad q$
$t = 2$	$\max\{\frac{1}{8}\frac{1}{2}\frac{1}{2}, \frac{1}{6}\frac{3}{4}\frac{1}{2}\} = \frac{3}{48} \quad p$	$\max\{\frac{1}{8}\frac{1}{2}\frac{2}{3}, \frac{1}{6}\frac{1}{4}\frac{2}{3}\} = \frac{1}{24} \quad q$
$t = 1$	$\frac{1}{2}\frac{1}{2}\frac{1}{2} = \frac{1}{8} \quad q$	$\frac{1}{2}\frac{1}{2}\frac{2}{3} = \frac{1}{6} \quad q$
$t = 0$	$\frac{1}{2} \quad q$	$0 \quad p$

Note that the two sequences of states, $qqpq$ and $qpqq$, are tied for the most likely sequences of states. ■

Determining the underlying hidden Markov model

Given an n -state HMM, how do we adjust the transition probabilities and output probabilities to maximize the probability of an output sequence $O_1O_2\cdots O_T$? The assumption is that T is much larger than n . There is no known computationally efficient method for solving this problem. However, there are iterative techniques that converge to a local optimum.

Let a_{ij} be the transition probability from state i to state j and let $b_j(O_k)$ be the probability of output O_k given that the HMM is in state j . Given estimates for the HMM parameters, a_{ij} and b_j , and the output sequence O , we can improve the estimates by calculating for each unit of time the probability that the HMM goes from state i to state j and outputs the symbol O_k .

a_{ij}	transition probability from state i to state j
$b_j(O_{t+1})$	probability of O_{t+1} given that the HMM is in state j at time $t + 1$
$\alpha_t(i)$	probability of seeing $O_0O_1\cdots O_t$ and ending in state i at time t
$\beta_{t+1}(j)$	probability of seeing the tail of the sequence $O_{t+2}O_{t+3}\cdots O_T$ given state j at time $t + 1$
$\delta(i, j)$	probability of going from state i to state j at time t given the sequence of outputs O
$s_t(i)$	probability of being in state i at time t given the sequence of outputs O
$p(O)$	probability of output sequence O

Given estimates for the HMM parameters, a_{ij} and b_j , and the output sequence O , the probability $\delta_t(i, j)$ of going from state i to state j at time t is given by the probability of producing the output sequence O and going from state i to state j at time t divided by the probability of producing the output sequence O .

$$\delta_t(i, j) = \frac{a_t(i)a_{ij}b_j(O_{t+1})\beta_{t+1}(j)}{p(O)}$$

The probability $p(O)$ is the sum over all pairs of states i and j of the numerator in the above formula for $\delta_t(i, j)$. That is,

$$p(O) = \sum_i \sum_j \alpha_t(j)a_{ij}b_j(O_{t+1})\beta_{t+1}(j).$$

The probability of being in state i at time t is given by

$$s_t(i) = \sum_{j=1}^n \delta_t(i, j).$$

Note that $\delta_t(i, j)$ is the probability of being in state i at time t given $O_0O_1O_2 \cdots O_t$ but it is not the probability of being in state i at time t given O since it does not take into account the remainder of the sequence O . Summing $s_t(i)$ over all time periods gives the expected number of times state i is visited and the sum of $\delta_t(i, j)$ over all time periods gives the expected number of times edge i to j is traversed.

Given estimates of the HMM parameters $a_{i,j}$ and $b_j(O_k)$, we can calculate by the above formulas estimates for

1. $\sum_{i=1}^{T-1} s_t(i)$, the expected number of times state i is visited and departed from
2. $\sum_{i=1}^{T-1} \delta_t(i, j)$, the expected number of transitions from state i to state j

Using these estimates we can obtain new estimates of the HMM parameters

$$\overline{a_{ij}} = \frac{\text{expected number of transitions from state } i \text{ to state } j}{\text{expected number of transitions out of state } i} = \frac{\sum_{t=1}^{T-1} \delta_t(i, j)}{\sum_{t=1}^{T-1} s_t(i)}$$

$$\overline{b_j}(O_k) = \frac{\text{expected number of times in state } j \text{ observing symbol } O_k}{\text{expected number of times in state } j} = \frac{\sum_{t=1}^{T-1} s_t(j)}{\sum_{t=1}^{T-1} s_t(j)} \quad \text{subject to } O_t = O_k$$

By iterating the above formulas we can arrive at a local optimum for the HMM parameters $a_{i,j}$ and $b_j(O_k)$.

9.3 Graphical Models, and Belief Propagation

A graphical model is a compact representation of a function of n variables $x_1, x_2, \dots, x_n$. It consists of a graph, directed or undirected, whose vertices correspond to variables that take on values from some set. In this chapter, we consider the case where the function is a probability distribution and the set of values the variables take on is finite, although graphical models are often used to represent probability distributions with continuous variables. The edges of the graph represent relationships or constraints between the variables.

The directed model represents a joint probability distribution that factors into a product of conditional probabilities.

$$p(x_1, x_2, \dots, x_n) = \prod_{i=1}^n p(x_i | \text{parents of } x_i)$$

It is assumed that the directed graph is acyclic. The directed graphical model is called a *Bayesian* or *belief network* and appears frequently in the artificial intelligence and the statistics literature.

The undirected graph model, called a Markov random field, can also represent a joint probability distribution of the random variables at its vertices. In many applications the Markov random field represents a function of the variables at the vertices which is to be optimized by choosing values for the variables.

A third model called the factor model is akin to the Markov random field, but here the dependency sets have a different structure. In the following sections we describe all these models in more detail.

9.4 Bayesian or Belief Networks

A Bayesian network is a directed acyclic graph where vertices correspond to variables and a directed edge from y to x represents a conditional probability $p(x|y)$. If a vertex x has edges into it from $y_1, y_2, \dots, y_k$, then the conditional probability is $p(x | y_1, y_2, \dots, y_k)$. The variable at a vertex with no in edges has an unconditional probability distribution. If the value of a variable at some vertex is known, then the variable is called *evidence*. An important property of a Bayesian network is that the joint probability is given by the product over all nodes of the conditional probability of the node conditioned on all its immediate predecessors.

In the example of Fig. 9.1, a patient is ill and sees a doctor. The doctor ascertains the symptoms of the patient and the possible causes such as whether the patient was in contact with farm animals, whether he had eaten certain foods, or whether the patient has an hereditary predisposition to any diseases. Using the above Bayesian network where the variables are true or false, the doctor may wish to determine one of two things. What

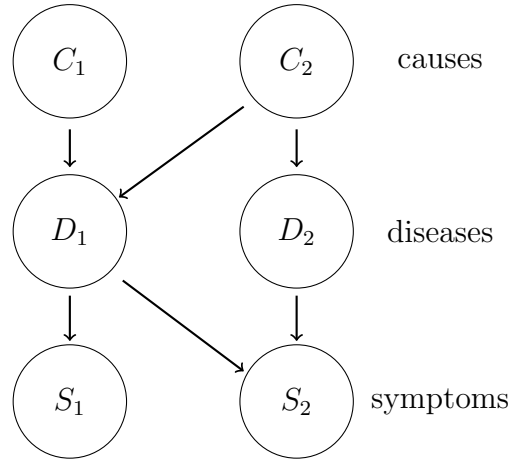


Figure 9.1: A Bayesian network

is the marginal probability of a given disease or what is the most likely set of diseases. In determining the most likely set of diseases, we are given a T or F assignment to the causes and symptoms and ask what assignment of T or F to the diseases maximizes the joint probability. This latter problem is called the maximum a posteriori probability (MAP).

Given the conditional probabilities and the probabilities $p(C_1)$ and $p(C_2)$ in Example 9.1, the joint probability $p(C_1, C_2, D_1, \dots)$ can be computed easily for any combination of values of $C_1, C_2, D_1, \dots$. However, we might wish to find that value of the variables of highest probability (MAP) or we might want one of the marginal probabilities $p(D_1)$ or $p(D_2)$. The obvious algorithms for these two problems require evaluating the probability $p(C_1, C_2, D_1, \dots)$ over exponentially many input values or summing the probability $p(C_1, C_2, D_1, \dots)$ over exponentially many values of the variables other than those for which we want the marginal probability. In certain situations, when the joint probability distribution can be expressed as a product of factors, a belief propagation algorithm can solve the maximum a posteriori problem or compute all marginal probabilities quickly.

9.5 Markov Random Fields

The Markov random field model arose first in statistical mechanics where it was called the Ising model. It is instructive to start with a description of it. The simplest version of the Ising model consists of n particles arranged in a rectangular $\sqrt{n} \times \sqrt{n}$ grid. Each particle can have a spin that is denoted ± 1 . The energy of the whole system depends on interactions between pairs of neighboring particles. Let x_i be the spin, ± 1 , of the i^{th} particle. Denote by $i \sim j$ the relation that i and j are adjacent in the grid. In the Ising model, the energy of the system is given by

$$f(x_1, x_2, \dots, x_n) = \exp \left(c \sum_{i \sim j} |x_i - x_j| \right).$$

c is a constant that can be positive or negative. If $c < 0$, then energy is lower if many adjacent pairs have opposite spins and if $c > 0$ the reverse holds. The model was first used to model probabilities of spin configurations. The hypothesis was that for each $\{x_1, x_2, \dots, x_n\}$ in $\{-1, +1\}^n$, the energy of the configuration with these spins is proportional to $f(x_1, x_2, \dots, x_n)$.

In most computer science settings, such functions are mainly used as objective functions that are to be optimized subject to some constraints. The problem is to find the minimum energy set of spins under some constraints on the spins. Usually the constraints just specify the spins of some particles. Note that when $c > 0$, this is the problem of minimizing $\sum_{i \sim j} |x_i - x_j|$ subject to the constraints. The objective function is convex and so this can be done efficiently. If $c < 0$, however, we need to minimize a concave function for which there is no known efficient algorithm. The minimization of a concave function in general is NP-hard.

A second important motivation comes from the area of vision. It has to do with reconstructing images. Suppose we are given observations of the intensity of light at individual pixels, $x_1, x_2, \dots, x_n$ and wish to compute the true values, the true intensities, of these variables $y_1, y_2, \dots, y_n$. There may be two sets of constraints, the first stipulating that the y_i must be close to the corresponding x_i and the second, a term correcting possible observation errors, stipulating that y_i must be close to the values of y_j for $j \sim i$. This can be formulated as

$$\text{Minimize } \sum_i |x_i - y_i| + \sum_{i \sim j} |y_i - y_j|,$$

where the values of x_i are constrained to be the observed values. The objective function is convex and polynomial time minimization algorithms exist. Other objective functions using say sum of squares instead of sum of absolute values can be used and there are polynomial time algorithms as long as the function to be minimized is convex.

More generally, the correction term may depend on all grid points within distance two of each point rather than just immediate neighbors. Even more generally, we may have n variables $y_1, y_2, \dots, y_n$ with the value of some already specified and subsets $S_1, S_2, \dots, S_m$ of these variables constrained in some way. The constraints are accumulated into one objective function which is a product of functions $f_1, f_2, \dots, f_m$, where function f_i is evaluated on the variables in subset S_i . The problem is to minimize $\prod_{i=1}^m f_i(y_j, j \in S_i)$ subject to constrained values. Note that the vision example had a sum instead of a product, but by taking exponentials we can turn the sum into a product as in the Ising model.

In general, the f_i are not convex; indeed they may be discrete. So the minimization cannot be carried out by a known polynomial time algorithm. The most used forms of the Markov random field involve S_i which are cliques of a graph. So we make the following definition.

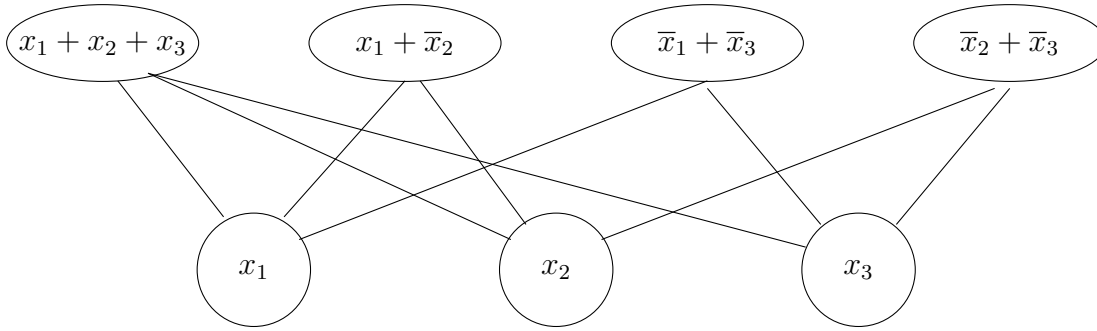


Figure 9.2: The factor graph for the function $f(x_1, x_2, x_3) = (x_1 + x_2 + x_3)(x_1 + x_2)(x_1 + x_3)(x_2 + x_3)$.

A *Markov Random Field* consists of an undirected graph and an associated function that factorizes into functions associated with the cliques of the graph. The special case when all the factors correspond to cliques of size one or two is of interest.

9.6 Factor Graphs

Factor graphs arise when we have a function f of a variables $\mathbf{x} = (x_1, x_2, \dots, x_n)$ that can be expressed as $f(\mathbf{x}) = \prod_{\alpha} f_{\alpha}(x_{\alpha})$ where each factor depends only on some small number of variables x_{α} . The difference from Markov random fields is that the variables corresponding to factors do not necessarily form a clique. Associate a bipartite graph where one set of vertices correspond to the factors and the other set to the variables. Place an edge between a variable and a factor if the factor contains that variable. See Figure 9.2

9.7 Tree Algorithms

Let $f(\mathbf{x})$ be a function that is a product of factors. When the factor graph is a tree there are efficient algorithms for solving certain problems. With slight modifications, the algorithms presented can also solve problems where the function is the sum of terms rather than a product of factors.

The first problem is called *marginalization* and involves evaluating the sum of f over all variables except one. In the case where f is a probability distribution the algorithm computes the marginal probabilities and thus the word marginalization. The second problem involves computing the assignment to the variables that maximizes the function f . When f is a probability distribution, this problem is the maximum a posteriori probability or MAP problem.

If the factor graph is a tree, then there exists an efficient algorithm for solving these problems. Note that there are four problems: the function f is either a product or a sum

and we are either marginalizing or finding the maximizing assignment to the variables. All four problems are solved by essentially the same algorithm and we present the algorithm for the marginalization problem when f is a product. Assume we want to “sum out” all the variables except x_1 , so we will be left with a function of x_1 .

We call the variable node associated with the variable x_i node x_i . First, make the node x_1 the root of the tree. It will be useful to think of the algorithm first as a recursive algorithm and then unravel the recursion. We want to compute the product of all factors occurring in the sub-tree rooted at the root with all variables except the root-variable summed out. Let g_i be the product of all factors occurring in the sub-tree rooted at node x_i with all variables occurring in the subtree except x_i summed out. Since this is a tree, x_1 will not reoccur anywhere except the root. Now, the grandchildren of the root are variable nodes and suppose for recursion, each grandchild x_i of the root, has already computed its g_i . It is easy to see that we can compute g_1 by the following.

Each grandchild x_i of the root passes its g_i to its parent, which is a factor node. Each child of x_1 collects all its children’s g_i , multiplies them together with its own factor and sends the product to the root. The root multiplies all the products it gets from its children and sums out all variables except its own variable, namely here x_1 .

Unraveling the recursion is also simple, with the convention that a leaf node just receives 1, product of an empty set of factors, from its children. Each node waits until it receives a message from each of its children. After that, if the node is a variable node, it computes the product of all incoming messages, and sums this product function over all assignments to the variables except for the variable of the node. Then, it sends the resulting function of one variable out along the edge to its parent. If the node is a factor node, it computes the product of its factor function along with incoming messages from all the children and sends the resulting function out along the edge to its parent.

The reader should prove that the following invariant holds assuming the graph is a tree:

Invariant The message passed by each variable node to its parent is the product of all factors in the subtree under the node with all variables in the subtree except its own summed out.

Consider the following example where

$$f = x_1 (x_1 + x_2 + x_3) (x_3 + x_4 + x_5) x_4 x_5$$

and the variables take on values 0 or 1. Consider marginalizing f by computing

$$f(x_1) = \sum_{x_2 x_3 x_4 x_5} x_1 (x_1 + x_2 + x_3) (x_3 + x_4 + x_5) x_4 x_5,$$

In this case the factor graph is a tree as shown in Figure 9.3. The factor graph as a

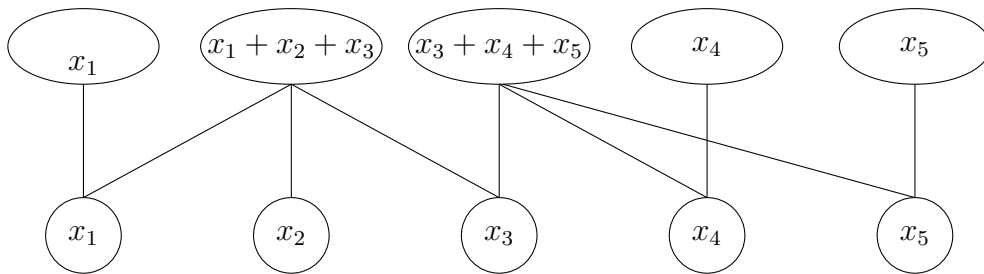


Figure 9.3: The factor graph for the function $f = x_1 (x_1 + x_2 + x_3) (x_3 + x_4 + x_5) x_4 x_5$.

rooted tree and the messages passed by each node to its parent are shown in Figure 9.4. If instead of computing marginal's, one wanted the variable assignment that maximizes the function f , one would modify the above procedure by replacing the summation by a maximization operation. Obvious modifications handle the situation where $f(\mathbf{x})$ is a sum of products.

$$f(\mathbf{x}) = \sum_{x_1, \dots, x_n} f(\mathbf{x})$$

9.8 Message Passing in general Graphs

The simple message passing algorithm in the last section gives us the one variable function of x_1 when we sum out all the other variables. For a general graph that is not a tree, we formulate an extension of that algorithm. But unlike the case of trees, there is no proof that the algorithm will converge and even if it does, there is no guarantee that the limit is the marginal probability. This has not prevented its usefulness in some applications.

First, lets ask a more general question, just for trees. Suppose we want to compute for each i the one variable function of x_i when we sum out all variables $x_j, j \neq i$. Do we have to repeat what we did for x_1 once for each x_i ? Luckily, the answer is no. It will suffice to do a second pass from the root to the leaves of essentially the same message passing algorithm to get all the answers. Recall that in the first pass, each edge of the tree has sent a message “up”, from the child to the parent. In the second pass, each edge will send a message from the parent to the child. We start with the root and work downwards for this pass. Each node waits until its parent has sent it a message before sending messages to each of its children. The rules for messages are:

Rule 1 The message from a factor node v to a child x_i , which is the variable node x_i , is the product of all messages received by v in both passes from all nodes other than x_i times the factor at v itself.

Rule 2 The message from a variable node x_i to a child, a factor node, v is the product of all messages received by x_i in both passes from all nodes except v , with all variables

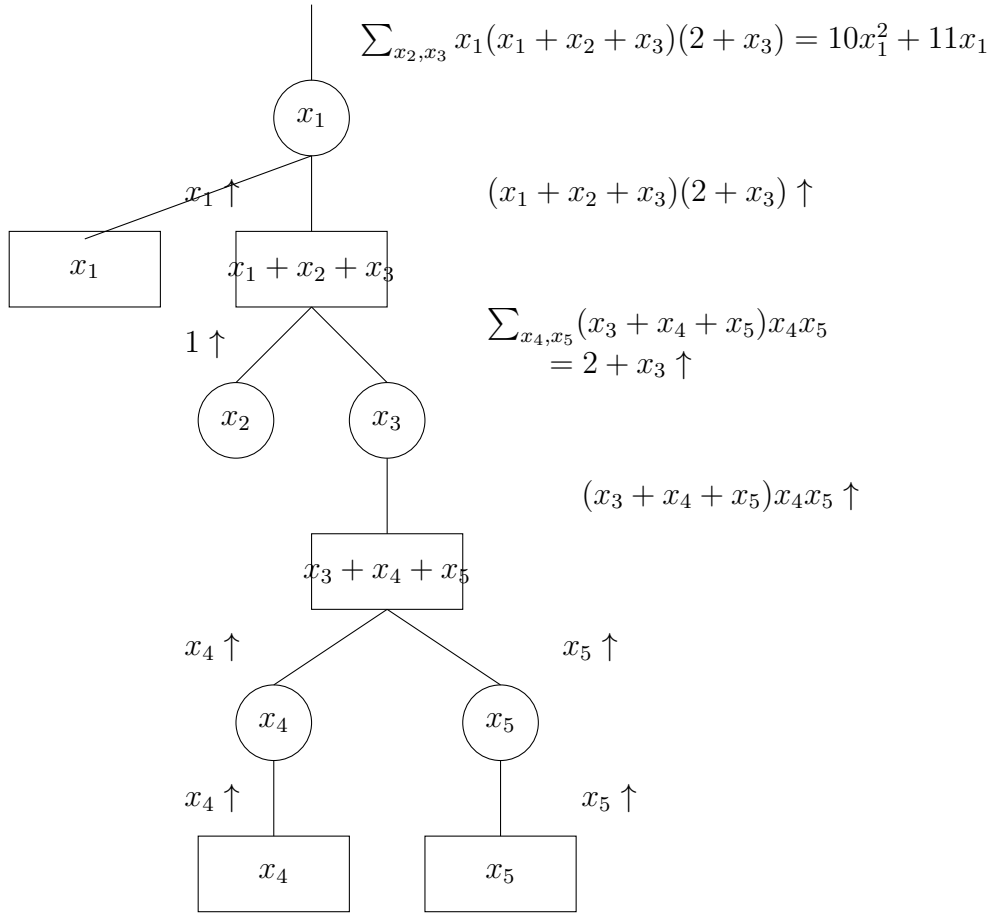


Figure 9.4: Messages.

except x_i summed out. The message is a function of x_i alone.

At termination, one can show when the graph is a tree that if we take the product of all messages received in both passes by a variable node x_i and sum out all variables except x_i in this product, what we get is precisely the entire function marginalized to x_i . We do not give the proof here. But the idea is simple. We know from the first pass that the product of the messages coming to a variable node x_i from its children is the product of all factors in the sub-tree rooted at x_i . In the second pass, we claim that the message from the parent v to x_i is the product of all factors which are not in the sub-tree rooted at x_i which one can show either directly or by induction working from the root downwards.

We can apply the same rules 1 and 2 to any general graph. We do not have child and parent relationships and it is not possible to have the two synchronous passes as before. The messages keep flowing and one hopes that after some time, the messages will stabilize, but nothing like that is proven. We state the algorithm for general graphs now:

Rule 1 At each time, each factor node v sends a message to each adjacent node x_i . The message is the product of all messages received by v at the previous step except for the one from x_i multiplied by the factor at v itself.

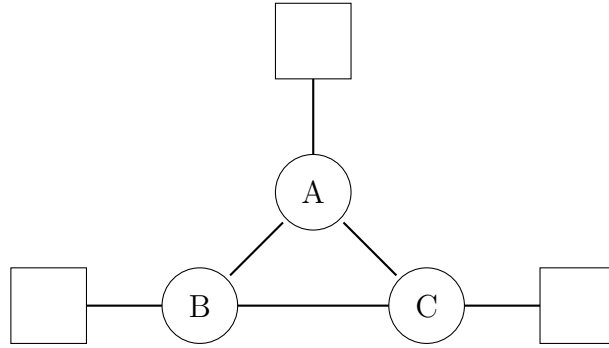
Rule 2 At each time, each variable node x_i sends a message to each adjacent node v . The message is the product of all messages received by x_i at the previous step except the one from v , with all variables except x_i summed out.

9.9 Graphs with a Single Cycle

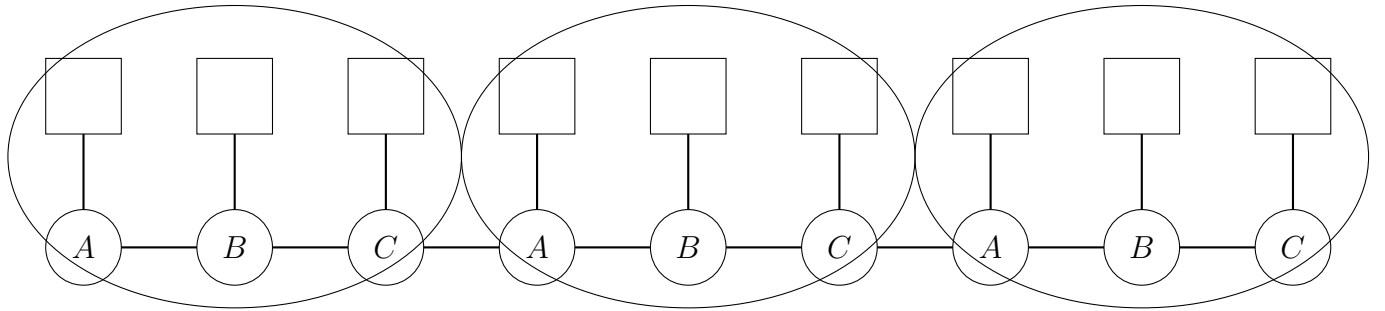
The message passing algorithm gives the correct answers on trees and on certain other graphs. One such situation is graphs with a single cycle which we treat here. We switch from the marginalization problem to the MAP problem as the proof of correctness is simpler for the MAP problem. Consider the network in the Figure 9.7a below with a single cycle. The message passing scheme will multiply count some evidence. The local evidence at A will get passed around the loop and will come back to A. Thus, A will count the local evidence multiple times. If all evidence is multiply counted in equal amounts, then there is a possibility that although the numerical values of the marginal probabilities (beliefs) are wrong, the algorithm still converges to the correct maximum a posteriori assignment.

Consider the unwrapped version of the graph in Figure 9.7b. The messages that the loopy version will eventually converge to, assuming convergence, are the same messages that occur in the unwrapped version provided that the nodes are sufficiently far in from the ends. The beliefs in the unwrapped version are correct for the unwrapped graph since it is a tree. The only question is, how similar are they to the true beliefs in the original network.

Write $p(A, B, C) = e^{\log p(A, B, C)} = e^{J(A, B, C)}$ where $J(A, B, C) = \log p(A, B, C)$. Then the probability for the unwrapped network is of the form $e^{kJ(A, B, C) + J'}$ where the J' is associated with vertices at the ends of the network where the beliefs have not yet stabilized and the $kJ(A, B, C)$ comes from k inner copies of the cycle where the beliefs have stabilized. Note that the last copy of J in the unwrapped network shares an edge with J' and that edge has an associated Ψ . Thus, changing a variable in J has an impact on the value of J' through that Ψ . Since the algorithm maximizes $J_k = kJ(A, B, C) + J'$ in the unwrapped network for all k , it must maximize $J(A, B, C)$. To see this, set the variables A, B, C, so that J_k is maximized. If $J(A, B, C)$ is not maximized, then change A, B, and C to maximize $J(A, B, C)$. This increases J_k by some quantity that is proportional to k . However, two of the variables that appear in copies of $J(A, B, C)$ also appear in J' and thus J' might decrease in value. As long as J' decreases by some finite amount, we can increase J_k by increasing k sufficiently. As long as all Ψ 's are nonzero, J' which is proportional to $\log \Psi$, can change by at most some finite amount. Hence, for a network with a single loop, assuming that the message passing algorithm converges, it converges



(a) A graph with a single cycle



(b) Segment of unrolled graph

Figure 9.5: Unwrapping a graph with a single cycle

to the maximum a posteriori assignment.

9.10 Belief Update in Networks with a Single Loop

In the previous section, we showed that when the message passing algorithm converges, it correctly solves the MAP problem for graphs with a single loop. The message passing algorithm can also be used to obtain the correct answer for the marginalization problem. Consider a network consisting of a single loop with variables $x_1, x_2, \dots, x_n$ and evidence $y_1, y_2, \dots, y_n$ as shown in Figure 9.6. The x_i and y_i can be represented by vectors having a component for each value x_i can take on. To simplify the discussion assume the x_i take on values $1, 2, \dots, m$.

Let m_i be the message sent from vertex i to vertex $i + 1 \bmod n$. At vertex $i + 1$ each component of the message m_i is multiplied by the evidence y_{i+1} and the constraint function Ψ . This is done by forming a diagonal matrix D_{i+1} where the diagonal elements are the evidence and then forming a matrix M_i whose rs^{th} element is $\Psi(x_{i+1} = r, x_i = s)$. The message m_{i+1} is $M_i D_{i+1} m_i$. Multiplication by the diagonal matrix D_{i+1} multiplies

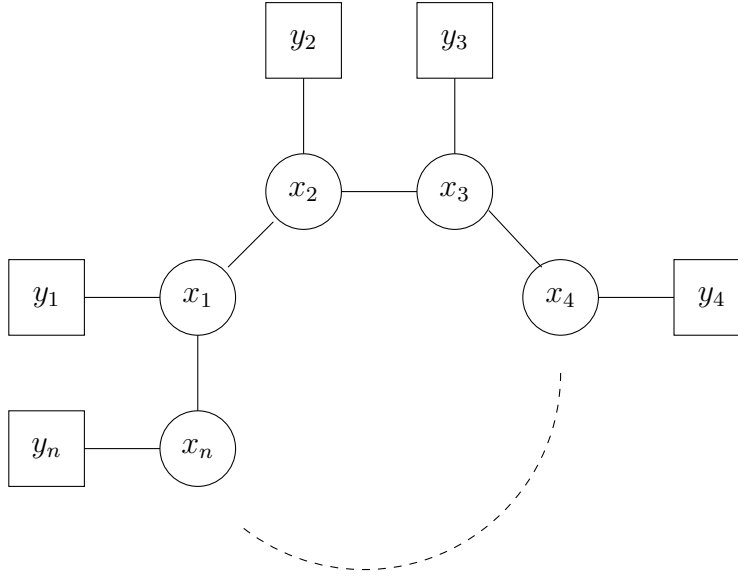


Figure 9.6: A Markov random field with a single loop.

the components of the message m_i by the associated evidence. Multiplication by the matrix M_i multiplies each component of the vector by the appropriate value of Ψ and sums over the values producing the vector which is the message m_{i+1} . Once the message has travelled around the loop, the new message m'_1 is given by

$$m'_1 = M_n D_1 M_{n-1} D_n \cdots M_2 D_3 M_1 D_2 m_1$$

Let $M = M_n D_1 M_{n-1} D_n \cdots M_2 D_3 M_1 D_2$. Assuming that M 's principle eigenvalue is unique, the message passing will converge to the principle vector of M . The rate of convergences depends on the ratio of the first and second eigenvalues.

An argument analogous to the above concerning the messages going clockwise around the loop applies to messages moving counter clockwise around the loop. To obtain the estimate of the marginal probability $p(x_1)$, one multiplies component wise the two messages arriving at x_1 along with the evidence y_1 . This estimate does not give the true marginal probability but the true marginal probability can be computed from the estimate and the rate of convergences by linear algebra.

9.11 Maximum Weight Matching

We have seen that the belief propagation algorithm converges to the correct solution in trees and graphs with a single cycle. It also correctly converges for a number of problems. Here we give one example, the maximum weight matching problem where there is a unique solution.

We apply the belief propagation algorithm to find the maximal weight matching (MWM) in a complete bipartite graph. If the MWM in the bipartite graph is unique, then the belief propagation algorithm will converge to it.

Let $G = (V_1, V_2, E)$ be a complete bipartite graph where $V_1 = \{a_1, \dots, a_n\}$, $V_2 = \{b_1, \dots, b_n\}$, and $(a_i, b_j) \in E$, $1 \leq i, j \leq n$. Let $\pi = \{\pi(1), \dots, \pi(n)\}$ be a permutation of $\{1, \dots, n\}$. The collection of edges $\{(a_1, b_{\pi(1)}), \dots, (a_n, b_{\pi(n)})\}$ is called a *matching* which is denoted by π . Let w_{ij} be the weight associated with the edge (a_i, b_j) . The weight of the matching π is $w_\pi = \sum_{i=1}^n w_{i\pi(i)}$. The maximum weight matching π^* is $\pi^* = \arg \max_{\pi} w_\pi$.

The first step is to create a factor graph corresponding to the MWM problem. Each edge of the bipartite graph is represented by a variable c_{ij} which takes on the values zero or one. The value one means that the edge is present in the matching, the value zero means that the edge is not present in the matching. A set of constraints is used to force the set of edges to be a matching. The constraints are of the form $\sum_j c_{ij} = 1$ and $\sum_i c_{ij} = 1$. Any assignment of 0,1 to the variables c_{ij} that satisfies all of the constraints defines a matching. In addition, we have constraints for the weights of the edges.

We now construct a factor graph, a portion of which is shown in Fig. 9.10. Associated with the factor graph is a function $f(c_{11}, c_{12}, \dots)$ consisting of a set of terms for each c_{ij} enforcing the constraints and summing the weights of the edges of the matching. The terms for c_{12} are

$$-\lambda \left| \left(\sum_i c_{i2} \right) - 1 \right| - \lambda \left| \left(\sum_j c_{1j} \right) - 1 \right| + w_{12} c_{12}$$

where λ is a large positive number used to enforce the constraints when we maximize the function. Finding the values of $c_{11}, c_{12}, \dots$ that maximize f finds the maximum weighted matching for the bipartite graph.

If the factor graph was a tree, then the message from a variable node x to its parent is a message $g(x)$ that gives the maximum value for the sub tree for each value of x . To compute $g(x)$, one sums all messages into the node x . For a constraint node, one sums all messages from sub trees and maximizes the sum over all variables except the variable of the parent node subject to the constraint. The message from a variable x consists of two pieces of information, the value $p(x=0)$ and the value $p(x=1)$. This information can be encoded into a linear function of x .

$$[p(x=1) - p(x=0)]x + p(x=0)$$

Thus, the messages are of the form $ax + b$. To determine the MAP value of x once the algorithm converges, sum all messages into x and take the maximum over $x=1$ and $x=0$.

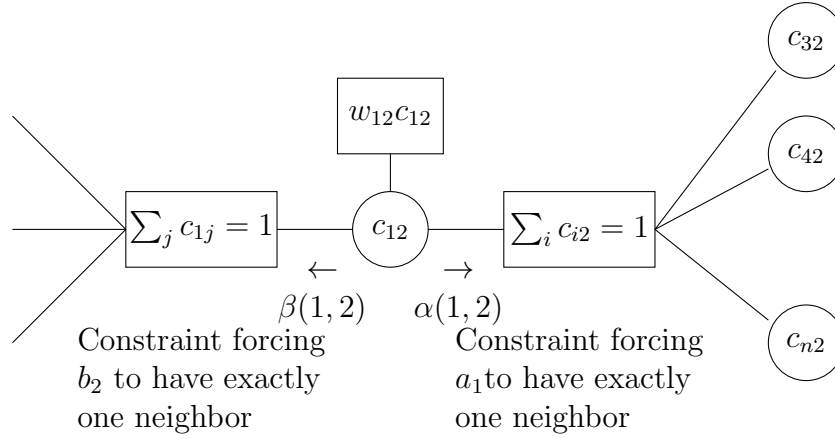


Figure 9.7: Portion of factor graph for the maximum weight matching problem.

to determine the value for x . Since the arg maximum of a linear form $ax+b$ depends only on whether a is positive or negative and since maximizing the output of a constraint depends only on the coefficient of the variable, we can send messages consisting of just the variable coefficient.

To calculate the message to c_{12} from the constraint that node b_2 has exactly one neighbor, add all the messages that flow into the constraint node from the c_{i2} , $i \neq 1$ nodes and maximize subject to the constraint that exactly one variable has value one. If $c_{12} = 0$, then one of c_{i2} , $i \neq 1$, will have value one and the message is $\max_{i \neq 1} \alpha(i, 2)$. If $c_{12} = 1$, then the message is zero. Thus, we get

$$-\max_{i \neq 1} \alpha(i, 2) x + \max_{i \neq 1} \alpha(i, 2)$$

and send the coefficient $-\max_{i \neq 1} \alpha(i, 2)$. This means that the message from c_{12} to the other constraint node is $\beta(1, 2) = w_{12} - \max_{i \neq 1} \alpha(i, 2)$.

The alpha message is calculated in a similar fashion. If $c_{12} = 0$, then one of c_{1j} will have value one and the message is $\max_{j \neq 1} \beta(1, j)$. If $c_{12} = 1$, then the message is zero. Thus, the coefficient $-\max_{j \neq 1} \alpha(1, j)$ is sent. This means that $\alpha(1, 2) = w_{12} - \max_{j \neq 1} \alpha(1, j)$.

To prove convergence, we enroll the constraint graph to form a tree with a constraint node as the root. In the enrolled graph a variable node such as c_{12} will appear a number of times which depends on how deep a tree is built. Each occurrence of a variable such as c_{12} is deemed to be a distinct variable.

Lemma 9.3 *If the tree obtained by unrolling the graph is of depth k , then the messages to the root are the same as the messages in the constraint graph after k -iterations.*

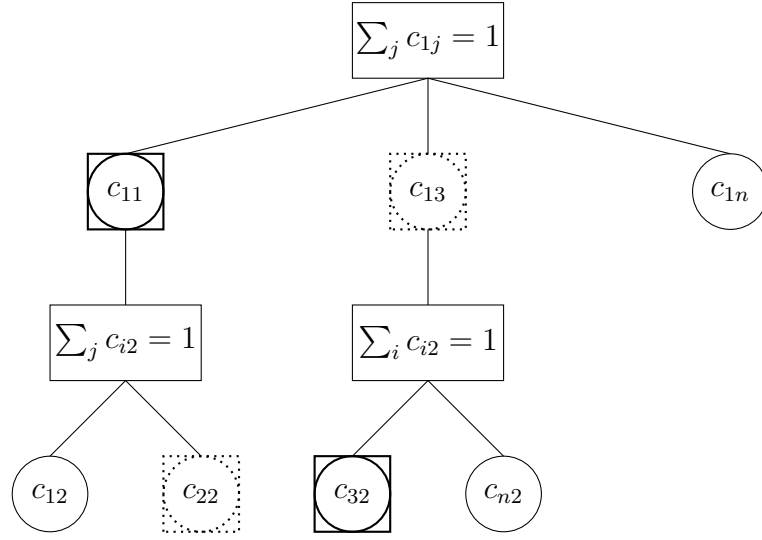


Figure 9.8: Tree for MWM problem.

Proof: Straight forward. ■

Define a matching in the tree to be a set of vertices so that there is exactly one variable node of the match adjacent to each constraint. Let Λ denote the vertices of the matching. Heavy circles represent the nodes of the above tree that are in the matching Λ .

Let Π be the vertices corresponding to maximum weight matching edges in the bipartite graph. Recall that vertices in the above tree correspond to edges in the bipartite graph. The vertices of Π are denoted by dotted circles in the above tree.

Consider a set of trees where each tree has a root that corresponds to one of the constraints. If the constraint at each root is satisfied by the edge of the MWM, then we have found the MWM. Suppose that the matching at the root in one of the trees disagrees with the MWM. Then there is an alternating path of vertices of length $2k$ consisting of vertices corresponding to edges in Π and edges in Λ . Map this path onto the bipartite graph. In the bipartite graph the path will consist of a number of cycles plus a simple path. If k is large enough there will be a large number of cycles since no cycle can be of length more than $2n$. Let m be the number of cycles. Then $m \geq \frac{2k}{2n} = \frac{k}{n}$.

Let π^* be the MWM in the bipartite graph. Take one of the cycles and use it as an alternating path to convert the MWM to another matching. Assuming that the MWM is unique and that the next closest matching is ε less, $W_{\pi^*} - W_{\pi} > \varepsilon$ where π is the new matching.

Consider the tree matching. Modify the tree matching by using the alternating path of all cycles and the left over simple path. The simple path is converted to a cycle by

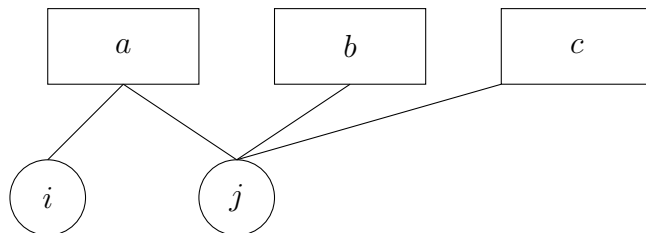


Figure 9.9: warning propagation

adding two edges. The cost of the two edges is at most $2w^*$ where w^* is the weight of the maximum weight edge. Each time we modify Λ by an alternating cycle, we increase the cost of the matching by at least ε . When we modify Λ by the left over simple path, we increase the cost of the tree matching by $\varepsilon - 2w^*$ since the two edges that were used to create a cycle in the bipartite graph are not used. Thus

$$\text{weight of } \Lambda - \text{weight of } \Lambda' \geq \frac{k}{n}\varepsilon - 2w^*$$

which must be negative since Λ' is optimal for the tree. However, if k is large enough this becomes positive, an impossibility since Λ' is the best possible. Since we have a tree, there can be no cycles, as messages are passed up the tree, each sub tree is optimal and hence the total tree is optimal. Thus the message passing algorithm must find the maximum weight matching in the weighted complete bipartite graph assuming that the maximum weight matching is unique. Note that applying one of the cycles that makes up the alternating path decreased the bipartite graph match but increases the value of the tree. However, it does not give a higher tree matching, which is not possible since we already have the maximum tree matching. The reason for this is that the application of a single cycle does not result in a valid tree matching. One must apply the entire alternating path to go from one matching to another.

9.12 Warning Propagation

Significant progress has been made using methods similar to belief propagation in finding satisfying assignments for 3-CNF formulas. Thus, we include a section on a version of belief propagation, called warning propagation, that is quite effective in finding assignments. Consider a factor graph for a SAT problem. Index the variables by i, j , and k and the factors by a, b , and c . Factor a sends a message m_{ai} to each variable i that appears in the factor a called a warning. The warning is 0 or 1 depending on whether or not factor a believes that the value assigned to i is required for a to be satisfied. A factor a determines the warning to send to variable i by examining all warnings received by other variables in factor a from factors containing them.

For each variable j , sum the warnings from factors containing j that warn j to take value T and subtract the warnings that warn j to take value F. If the difference says that

j should take value T or F and this value for variable j does not satisfy a , and this is true for all j , then a sends a warning to i that the value of variable i is critical for factor a .

Start the warning propagation algorithm by assigning 1 to a warning with probability $1/2$. Iteratively update the warnings. If the warning propagation algorithm converges, then compute for each variable i the local field h_i and the contradiction number c_i . The local field h_i is the number of clauses containing the variable i that sent messages that i should take value T minus the number that sent messages that i should take value F. The contradiction number c_i is 1 if variable i gets conflicting warnings and 0 otherwise. If the factor graph is a tree, the warning propagation algorithm converges. If one of the warning messages is one, the problem is unsatisfiable; otherwise it is satisfiable.

9.13 Correlation Between Variables

In many situations one is interested in how the correlation between variables drops off with some measure of distance. Consider a factor graph for a 3-CNF formula. Measure the distance between two variables by the shortest path in the factor graph. One might ask the question if one variable is assigned the value true, what is the percentage of satisfying assignments in which the second variable also is true. If the percentage is the same as when the first variable is assigned false, then we say that the two variables are uncorrelated. How difficult it is to solve a problem is likely to be related to how fast the correlation decreases with distance.

Another illustration of this concept is in counting the number of perfect matchings in a graph. One might ask what is the percentage of matching in which some edge is present and ask how correlated this percentage is with the presences or absence of edges at some distance d . One is interested in whether the correlation drops off with distance d . To illustrate this concept we consider the Ising model studied in physics.

The Ising or ferromagnetic model is a pairwise random Markov field. The underlying graph, usually a lattice, assigns a value of ± 1 , called spin, to the variable at each vertex. The probability (Gibbs measure) of a given configuration of spins is proportional to $\exp(\beta \sum_{(i,j) \in E} x_i x_j) = \prod_{(i,j) \in E} e^{\beta x_i x_j}$ where $x_i = \pm 1$ is the value associated with vertex i .

Thus

$$p(x_1, x_2, \dots, x_n) = \frac{1}{Z} \prod_{(i,j) \in E} \exp(\beta x_i x_j) = \frac{1}{Z} e^{\beta \sum_{(i,j) \in E} x_i x_j}$$

where Z is a normalization constant.

The value of the summation is simply the difference in the number of edges whose vertices have the same spin minus the number of edges whose vertices have opposite spin. The constant β is viewed as inverse temperature. High temperature corresponds to a low value of β . At low temperature, high β , adjacent vertices have identical spins whereas at

high temperature the spins of adjacent vertices are uncorrelated.

One question of interest is given the above probability distribution, what is the correlation between two variables, say x_i and x_j . To answer this question, we want to determine the $Prob(x_i = 1)$ as a function of $Prob(x_j = 1)$. If $Prob(x_i = 1) = \frac{1}{2}$ independent of the value of $Prob(x_j = 1)$, we say the values are uncorrelated.

Consider the special case where the graph G is a tree. In this case a phase transition occurs at $\beta_0 = \frac{1}{2} \ln \frac{d+1}{d-1}$ where d is the degree of the tree. For a sufficiently tall tree and for $\beta > \beta_0$, the probability that the root has value $+1$ is bounded away from $\frac{1}{2}$ and depends on whether the majority of leaves have value $+1$ or -1 . For $\beta < \beta_0$ the probability that the root has value $+1$ is $\frac{1}{2}$ independent of the values at the leaves of the tree.

Consider a height one tree of degree d . If i of the leaves have spin $+1$ and $d - i$ have spin -1 , then the probability of the root having spin $+1$ is proportional to

$$e^{i\beta - (d-i)\beta} = e^{(2i-d)\beta}.$$

If the probability of a leaf being $+1$ is p , then the probability of i leaves being $+1$ and $d - i$ being -1 is

$$\binom{d}{i} p^i (1-p)^{d-i}$$

Thus, the probability of the root being $+1$ is proportional to

$$A = \sum_{i=1}^d \binom{d}{i} p^i (1-p)^{d-i} e^{(2i-d)\beta} = e^{-d\beta} \sum_{i=1}^d \binom{d}{i} (pe^{2\beta})^i (1-p)^{d-i} = e^{-d\beta} [pe^{2\beta} + 1 - p]^d$$

and the probability of the root being -1 is proportional to

$$B = \sum_{i=1}^d \binom{d}{i} p^i (1-p)^{d-i} e^{-(2i-d)\beta} = e^{-d\beta} \sum_{i=1}^d \binom{d}{i} (p)^i [(1-p)e^{2\beta}]^{d-i} = e^{-d\beta} [p + (1-p)e^{2\beta}]^d.$$

The probability of the root being $+1$ is

$$q = \frac{A}{A+B} = \frac{[pe^{2\beta} + 1 - p]^d}{[pe^{2\beta} + 1 - p]^d + [p + (1-p)e^{2\beta}]^d} = \frac{C}{D}$$

where

$$C = [pe^{2\beta} + 1 - p]^d$$

and

$$D = [pe^{2\beta} + 1 - p]^d + [p + (1-p)e^{2\beta}]^d.$$

At high temperature the probability q of the root of this height one tree being 1 is $1/2$ independent of p . At low temperature q goes from low probability of 1 below $p=1/2$ to high probability of 1 above $p=1/2$.

Now consider a very tall tree. If the p is the probability that a root has value $+1$, we can iterate the formula for the height one tree and observe that at low temperature the probability of the root being one converges to some value. At high temperature, the probability of the root being one is $1/2$ independent of p . At the phase transition, the slope of q at $p=1/2$ is one.

Now the slope of the probability of the root being 1 with respect to the probability of a leaf being 1 in this height one tree is

$$\frac{\partial q}{\partial p} = \frac{D \frac{\partial C}{\partial p} - C \frac{\partial D}{\partial p}}{D^2}$$

Since the slope of the function $q(p)$ at $p=1/2$ when the phase transition occurs is one, we can solve $\frac{\partial q}{\partial p} = 1$ for the value of β where the phase transition occurs. First, we show that

$$\left. \frac{\partial D}{\partial p} \right|_{p=\frac{1}{2}} = 0.$$

$$\begin{aligned} D &= [pe^{2\beta} + 1 - p]^d + [p + (1 - p)e^{2\beta}]^d \\ \frac{\partial D}{\partial p} &= d [pe^{2\beta} + 1 - p]^{d-1} (e^{2\beta} - 1) + d [p + (1 - p)e^{2\beta}]^{d-1} (1 - e^{2\beta}) \\ \left. \frac{\partial D}{\partial p} \right|_{p=\frac{1}{2}} &= \frac{d}{2^{d-1}} [e^{2\beta} + 1]^{d-1} (e^{2\beta} - 1) + \frac{d}{2^{d-1}} [1 + e^{2\beta}]^{d-1} (1 - e^{2\beta}) = 0 \end{aligned}$$

Then

$$\begin{aligned} \left. \frac{\partial q}{\partial p} \right|_{p=\frac{1}{2}} &= \frac{D \frac{\partial C}{\partial p} - C \frac{\partial D}{\partial p}}{D^2} \Big|_{p=\frac{1}{2}} = \frac{\frac{\partial C}{\partial p}}{D} \Big|_{p=\frac{1}{2}} = \frac{d [pe^{2\beta} + 1 - p]^{d-1} (e^{2\beta} - 1)}{[pe^{2\beta} + 1 - p]^d + [p + (1 - p)e^{2\beta}]^d} \Big|_{p=\frac{1}{2}} \\ &= \frac{d [\frac{1}{2}e^{2\beta} + \frac{1}{2}]^{d-1} (e^{2\beta} - 1)}{[\frac{1}{2}e^{2\beta} + \frac{1}{2}]^d + [\frac{1}{2} + \frac{1}{2}e^{2\beta}]^d} = \frac{d (e^{2\beta} - 1)}{1 + e^{2\beta}} \end{aligned}$$

Setting

$$\frac{d (e^{2\beta} - 1)}{1 + e^{2\beta}} = 1$$

And solving for β yields

$$\begin{aligned} d (e^{2\beta} - 1) &= 1 + e^{2\beta} \\ e^{2\beta} &= \frac{d+1}{d-1} \\ \beta &= \frac{1}{2} \ln \frac{d+1}{d-1} \end{aligned}$$

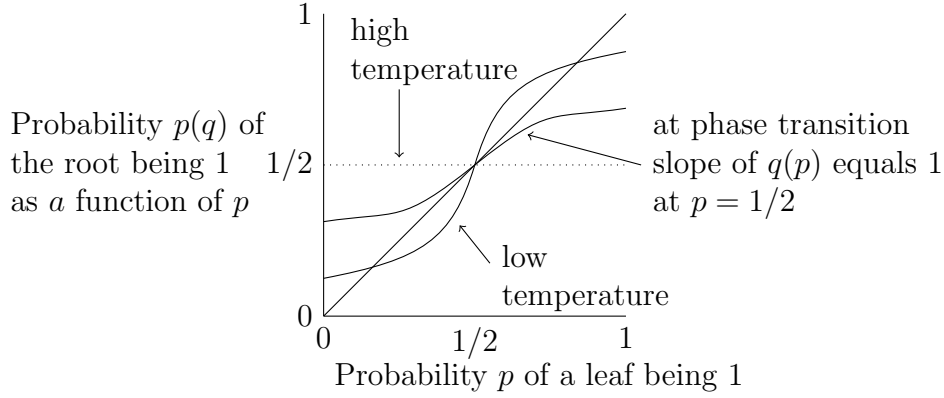


Figure 9.10: Shape of q as a function of p .

To complete the argument, we need to show that q is a monotonic function of p . To see this, write $q = \frac{1}{1 + \frac{B}{A}}$. A is a monotonically increasing function of p and B is monotonically decreasing. From this it follows that q is monotonically increasing.

In the iteration going from p to q , we do not get the true marginal probabilities at each level since we ignored the effect of the portion of the tree above. However, when we get to the root, we do get the true marginal for the root. To get the true marginal's for the interior nodes we need to send messages down from the root.

Note: The joint probability distribution for the tree is of the form $e^{\beta \sum_{(i,j) \in E} x_i x_j} = \prod_{(i,j) \in E} e^{\beta x_i x_j}$.

Suppose x_1 has value 1 with probability p . Then define a function φ , called evidence, such that

$$\begin{aligned} \varphi(x_1) &= \begin{cases} p & \text{for } x_1 = 1 \\ 1 - p & \text{for } x_1 = -1 \end{cases} \\ &= \left(p - \frac{1}{2}\right) x_1 + \frac{1}{2} \end{aligned}$$

and multiply the joint probability function by φ . Note, however, that the marginal probability of x_1 is not p . In fact, it may be further from p after multiplying the conditional probability function by the function φ .

9.14 Exercises

Exercise 9.1 Find a nonnegative factorization of the matrix

$$A = \begin{pmatrix} 4 & 6 & 5 \\ 1 & 2 & 3 \\ 7 & 10 & 7 \\ 6 & 8 & 4 \\ 6 & 10 & 11 \end{pmatrix}$$

Indicate the steps in your method and show the intermediate results.

Exercise 9.2 Find a nonnegative factorization of each of the following matrices.

$$(1) \begin{pmatrix} 10 & 9 & 15 & 14 & 13 \\ 2 & 1 & 3 & 3 & 1 \\ 8 & 7 & 13 & 11 & 11 \\ 7 & 5 & 11 & 10 & 7 \\ 5 & 5 & 11 & 6 & 11 \\ 1 & 1 & 3 & 1 & 3 \\ 2 & 2 & 2 & & 2 \end{pmatrix}$$

$$(2) \begin{pmatrix} 5 & 5 & 10 & 14 & 17 \\ 2 & 2 & 4 & 4 & 6 \\ 1 & 1 & 2 & 4 & 4 \\ 1 & 1 & 2 & 2 & 3 \\ 3 & 3 & 6 & 8 & 10 \\ 5 & 5 & 10 & 16 & 18 \\ 2 & 2 & 4 & 6 & 7 \end{pmatrix}$$

$$(3) \begin{pmatrix} 4 & 4 & 3 & 3 & 1 & 3 & 4 & 3 \\ 13 & 16 & 13 & 10 & 5 & 13 & 14 & 10 \\ 15 & 24 & 21 & 12 & 9 & 21 & 18 & 12 \\ 7 & 16 & 15 & 6 & 7 & 15 & 10 & 6 \\ 1 & 4 & 4 & 1 & 2 & 4 & 2 & 1 \\ 5 & 8 & 7 & 4 & 3 & 7 & 6 & 4 \\ 3 & 12 & 12 & 3 & 6 & 12 & 6 & 3 \end{pmatrix}$$

$$(4) \begin{pmatrix} 1 & 1 & 3 & 4 & 4 & 4 & 1 \\ 9 & 9 & 9 & 12 & 9 & 9 & 3 \\ 6 & 6 & 12 & 16 & 15 & 15 & 4 \\ 3 & 3 & 3 & 4 & 3 & 3 & 1 \end{pmatrix}$$

Exercise 9.3 Consider the matrix A that is the product of nonnegative matrices B and C .

$$\begin{pmatrix} 12 & 22 & 41 & 35 \\ 19 & 20 & 13 & 48 \\ 11 & 14 & 16 & 29 \\ 14 & 16 & 14 & 36 \end{pmatrix} = \begin{pmatrix} 10 & 1 \\ 1 & 9 \\ 3 & 4 \\ 2 & 6 \end{pmatrix} \begin{pmatrix} 1 & 2 & 4 & 3 \\ 2 & 2 & 1 & 5 \end{pmatrix}$$

Which rows of A are approximate positive linear combinations of other rows of A ?
Find an approximate nonnegative factorization of A

Exercise 9.4 What is the probability of heads occurring after a sufficiently long sequence of transitions in example 9.2?

Exercise 9.5 Find optimum parameters for a three state HMM and given output sequence. Note the HMM must a strong signature in the output sequence or we probably will not be able to find it. The following example may not be good for that reason.

	1	2	3		A	B
1	$\frac{1}{2}$	$\frac{1}{4}$	$\frac{1}{4}$	1	$\frac{3}{4}$	$\frac{1}{4}$
2	$\frac{1}{4}$	$\frac{1}{4}$	$\frac{1}{2}$	2	$\frac{1}{4}$	$\frac{3}{4}$
3	$\frac{1}{3}$	$\frac{1}{3}$	$\frac{1}{3}$	3	$\frac{1}{3}$	$\frac{2}{3}$

Exercise 9.6 In the Ising model for a tree of degree one, a chain of vertices, is there a phase transition where the correlation between the value at the root and the value at the leaves becomes independent? Work out mathematically what happens.

Exercise 9.7 Consider an n by n grid of vertices with a variable at each vertex. Define the single loop and tree neighborhood (SLT) of an assignment r to the variables to be the set of all assignments that differ from the assignment r by changing the value of a set of variables whose induced graph is a set of trees plus possibly a single loop. Prove that any assignment to the variables is in the SLT neighborhood of an assignment v where v is in the SLT neighborhood of the optimal assignment.

Exercise 9.8 What happens in Exercise 9.2 if the graph is an n -clique? What if the underlying graph is regular degree d ?

Exercise 9.9 Research question: How many solutions are required for 2-d grid so that every solution is in one of their SLT neighborhoods?

Exercise 9.10 For a Boolean function in CNF the marginal probability gives the number of satisfiable assignments with x_1 .

How does one obtain the number of satisfying assignments for a 2-CNF formula? Not completely related to first sentence.

References

“Factor graphs and the sum product algorithm”

Yair Weiss, “Belief propagation and revision in networks with loops” 1997

“On the optimality of solutions of the max-product belief propagation algorithm in arbitrary graphs” Yair Weiss and William T. Freeman

Brendan J. Frey and Delbert Dueck, “Clustering by passing messages between data points”, Science 315 Feb 16, 2007 pp972-976.

Bethe free energy, Kikuchi approximations, and belief propagation algorithms, Jonathan S. Yedidia, William T. Freeman, and Yair Weiss 2001

Jonathan S. Yedidia, William T. Freeman, and Yair Weiss, Understanding Belief Propagation and its Generalizations. 2002

Frank R. Kschischang, Brendan J. Frey, and Hans-Andrea Loeliger, “Factor Graphs and the Sum-Product Algorithm, IEEE Transactions on Information Theory Vol 47:2 Feb 2001.

M. Mezard, G. Parisi, R. Zecchina, “Analytic and Algorithmic Solution of Random Satisfiability Problems, Science Vol 297 no 5582 pp812-815, 2002

A. Braunstein, M. Mezard, and R. Zecchina, “Survey propagation: an algorithm for satisfiability” 2004

Dimitris Achlioptas and Federico Ricci-Tersenghi, “On the solution-space geometry of random constraint satisfaction problems”, STOC 2006

Bayati, Shah, and Sharma, “Max-product for maximum weight matching: convergence, correctness and LP duality. Good overview of field with references. Claims single cycle graphs rigorously analyzed

“Exact inference using the attenuated max-product algorithm”, Brendan J. Frey and Ralf Koetter July 31, 2000

Notes on warning propagation taken from A. Braunstein, M. Mezard, and R. Zecchina, “Survey propagation: an algorithm for satisfiability” 2004

10 Other Topics

10.1 Rankings

Ranking is important. We rank movies, restaurants, students, web pages, and many other items. Ranking has become a multi-billion dollar industry as organizations try to raise the position of their web pages in the display of web pages returned by search engines to relevant queries. Developing a method of ranking that is not manipulative is an important task.

A ranking is a complete ordering in the sense that for every pair of items a and b , either a is preferred to b or b is preferred to a . Furthermore, a ranking is transitive in that $a > b$ and $b > c$ implies $a > c$.

One problem of interest in ranking is that of combining many individual rankings into one global ranking. However, merging ranked lists is nontrivial as the following example illustrates.

Example: Suppose there are three individuals who rank items a , b , and c as illustrated in the following table.

individual	first item	second item	third item
1	a	b	c
2	b	c	a
3	c	a	b

Suppose our algorithm tried to rank the items by first comparing a to b and then comparing b to c . In comparing a to b , two of the three individuals prefer a to b and thus we conclude a is preferable to b . In comparing b to c , again two of the three individuals prefer b to c and we conclude that b is preferable to c . Now by transitivity one would expect that the individuals would prefer a to c , but such is not the case, only one of the individuals prefers a to c and thus c is preferable to a . We come to the illogical conclusion that a is preferable to b , b is preferable to c , but c is preferable to a . ■

Suppose there are a number of individuals or voters and a set of candidates to be ranked. Each voter produces a ranked list of the candidates. From the set of ranked lists can one construct a single ranking of the candidates? Assume the method of producing a global ranking is required to satisfy the following three axioms.

Nondictatorship – The algorithm cannot always simply select one individual's ranking.

Unanimity – If every individual prefers a to b , then the global ranking must prefer a to b .

Independent of irrelevant alternatives – If individuals modify their rankings but keep the order of a and b unchanged, then the global order of a and b should not change.

Arrow showed that no ranking algorithm exists satisfying the above axioms.

Theorem 10.1 (Arrow) *Any algorithm for creating a global ranking of three or more elements that satisfies unanimity and independence of irrelevant alternatives is a dictatorship.*

Proof: Let a , b , and c be distinct items. Consider a set of rankings in which each individual ranks b either first or last. Some individuals may rank b first and others may rank b last. For this set of rankings the global ranking must put b first or last. Suppose to the contrary that b is not first or last in the global ranking. Then there exist a and c where the global ranking puts $a > b$ and $b > c$. By transitivity, the global ranking puts $a > c$. Note that all individuals can move c above a without affecting the order of b and a or the order of b and c since b was first or last on each list. Thus, by independence of irrelevant alternatives, the global ranking would continue to rank $a > b$ and $b > c$ even if all individuals moved c above a since that would not change the individuals relative order of a and b or the individuals relative order of b and c . But then by unanimity, the global ranking would need to put $c > a$, a contradiction. We conclude that the global ranking puts b first or last.

Consider a set of rankings in which every individual ranks b last. By unanimity, the global ranking must also rank b last. Let the individuals, one by one, move b from bottom to top leaving the other rankings in place. By unanimity, the global ranking must eventually move b from the bottom all the way to the top. When b first moves, it must move all the way to the top by the previous argument. Let v be the first individual whose change causes the global ranking of b to change.

We now argue that v is a dictator. First, we argue that v is a dictator for any pair ac not involving b . We will refer to three rankings of v (see Figure 10.1). The first ranking of v is the ranking prior to v moving b from the bottom to the top and the second is the ranking just after v has moved b to the top. Choose any pair ac where a is above c in v 's ranking. The third ranking of v is obtained by moving a above b so that $a > b > c$ in v 's ranking. By independence of irrelevant alternatives, the global ranking after v has switched to the third ranking puts $a > b$ since all individual ab votes are the same as just before v moved b to the top of his ranking. At that time the global ranking placed $a > b$. Similarly $b > c$ in the global ranking since all individual bc votes are the same as just after v moved b to the top causing b to move to the top in the global ranking. By transitivity the global ranking must put $a > c$ and thus the global ranking of a and c agrees with v .

Now all individuals except v can modify their rankings arbitrarily while leaving b in its extreme position and by independence of irrelevant alternatives, this does not affect the global ranking of $a > b$ or of $b > c$. Thus, by transitivity this does not affect the global

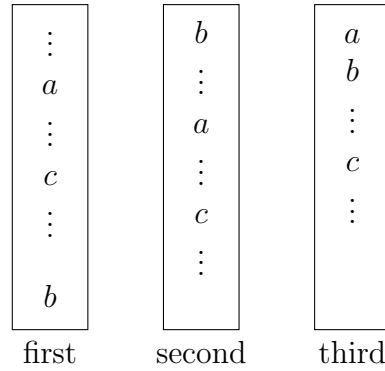


Figure 10.1: Rankings of v

ranking of a and c . Next, all individuals except v can move b to any position without affecting the global ranking of a and c .

At this point we have argued that independent of other individuals' rankings, the global ranking of a and c will agree with v 's ranking. Now v can change its ranking arbitrarily, provided it maintains the order of a and c , and by independence of irrelevant alternatives the global ranking of a and c will not change and hence will agree with v . Thus, we conclude that for all a and c , the global ranking agrees with v independent of the other rankings except for the placement of b . But other rankings can move b without changing the global order of other elements. Thus, v is a dictator for the ranking of any pair of elements not involving b .

Note that v changed the relative order of a and b in the global ranking when it moved b from the bottom to the top in the previous argument. We will use this in a moment.

The individual v is also a dictator over every pair ab . Repeat the construction showing that v is a dictator for every pair ac not involving b only this time place c at the bottom. There must be an individual v_c who is a dictator for any pair such as ab not involving c . Since both v and v_c can affect the global ranking of a and b independent of each other, it must be that v_c is actually v . Thus, the global ranking agrees with v no matter how the other voters modify their rankings. ■

10.2 Hare System for Voting

One voting system would be to have everyone vote for their favorite candidate. If some candidate receives a majority of votes, he or she is declared the winner. If no candidate receives a majority of votes, the candidate with the fewest votes is dropped from the slate and the process is repeated.

The Hare system implements this method by asking each voter to rank all the can-

didates. Then one counts how many voters ranked each candidate as number one. If no candidate receives a majority, the candidate with the fewest number one votes is dropped from each voters ranking. If the dropped candidate was number one on some voters list, then the number two candidate becomes that voter's number one choice. The process of counting the number one rankings is then repeated.

Although the Hare system is widely used it fails to satisfy Arrow's axioms as all voting systems must. Consider the following situation in which there are 21 voters that fall into four categories. Voters within a category rank individuals in the same order.

Category	Number of voters in category	Preference order
1	7	abcd
2	6	bacd
3	5	cbad
4	3	dcba

The Hare system would first eliminate d since d gets only three rank one votes. Then it would eliminate b since b gets only six rank one votes whereas a gets seven and c gets eight. At this point a is declared the winner since a has thirteen votes to c 's eight votes.

Now assume that Category 4 voters who prefer b to a move a up to first place. Then the election proceeds as follows. In round one, d is eliminated since it gets no rank 1 votes. Then c with five votes is eliminated and b is declared the winner with 11 votes. Note that by moving a up, category 4 voters were able to deny a the election and get b to win, whom they prefer over a .

10.3 Compressed Sensing and Sparse Vectors

If one has a discrete time sequence $\mathbf{x}$ of length n , the Nyquist theorem states that n coefficients in the frequency domain are needed to represent the signal $\mathbf{x}$. However, if the signal $\mathbf{x}$ has only s nonzero elements, even though one does not know which elements they are, one can recover the signal by randomly selecting a small subset of the coefficients in the frequency domain. It turns out that one can reconstruct sparse signals with far fewer samples than one might suspect and an area called compressed sampling has emerged with important applications.

Motivation

Let A be an $n \times d$ matrix with n much smaller than d whose elements are generated by independent Gaussian processes. Let $\mathbf{x}$ be a sparse d -dimensional vector with at most s non-zero coordinates, $s \ll d$. $\mathbf{x}$ is called the signal and A is the "measurement" matrix.

$$A x = b$$

Figure 10.2: $A\mathbf{x} = \mathbf{b}$ has a vector space of solutions but possibly only one sparse solution.

What we measure are the components of the n dimensional vector $A\mathbf{x}$. We ask if we can recover the signal $\mathbf{x}$ from measurements $A\mathbf{x}$, where, the number n of measurements is much smaller than the dimension d ? We have two advantages over an arbitrary system of linear equations. First, the solution $\mathbf{x}$ is known to be sparse and second we have the choice of the measurement matrix A .

While we do not describe the motivation for this problem in any detail, here it is in brief. In many applications, the signal is sparse in either the time domain or the frequency domain. For example, for images, it is often the case that in the frequency domain, very few frequencies have significant amplitude and if we zero out small amplitudes, we get a sparse signal. It is wasteful to measure each of the d components of the signal $\mathbf{x}$. Instead, we measure n linear combinations of components, the linear combinations form our A . In applications, we choose the matrix A . An usual choice is a matrix whose entries are independent zero mean, unit variance Gaussian random variables. Since we have no control over the signal, our system needs to recover any signal. We will show that n needs to depend essentially only on s , not on d .

10.3.1 Unique Reconstruction of a Sparse Vector

A vector is said to be s -sparse if it has at most s nonzero elements. Let $\mathbf{x}$ be an d -dimensional s -sparse vector with $s \ll d$. Consider solving $A\mathbf{x} = \mathbf{b}$ for $\mathbf{x}$ where A is an $n \times d$ matrix with $n < d$. The set of solutions to $A\mathbf{x} = \mathbf{b}$ is a subspace. However, if we restrict ourselves to sparse solutions, under certain conditions on A there is a unique s -sparse solution. Suppose that there were two s -sparse solutions, $\mathbf{x}_1$ and $\mathbf{x}_2$. Then $\mathbf{x}_1 - \mathbf{x}_2$ would be a $2s$ -sparse solution to the homogeneous system $A\mathbf{x} = \mathbf{0}$. A $2s$ -sparse solution to the homogeneous equation $A\mathbf{x} = \mathbf{0}$ requires that some $2s$ columns of A be linearly dependent. Unless A has $2s$ linearly dependent columns there can be only one s -sparse solution.

Now suppose $n \in \Omega(s^2)$ and we pick an $n \times d$ matrix A with random independent

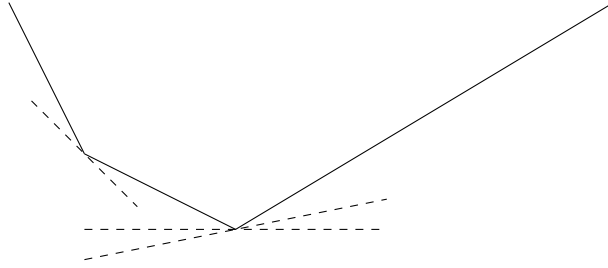


Figure 10.3: Some subgradients for a function that is not everywhere differentiable.

zero mean, unit variance Gaussian entries. Take any subset of $2s$ columns of A . Since we have already seen in Chapter 2 that each of these $2s$ vectors is likely to be essentially orthogonal to the space spanned by the previous vectors, the sub-matrix is unlikely to be singular. This is only an intuition, but it can be made rigorous.

To find a sparse solution to $A\mathbf{x} = \mathbf{b}$, one would like to minimize the zero norm $\|\mathbf{x}\|_0$ over $\{\mathbf{x} | A\mathbf{x} = \mathbf{b}\}$. The difficulty is that this is a computationally hard problem. We have techniques to minimize a convex function over a convex set. But $\|\mathbf{x}\|_0$ is not a convex function. With no further hypotheses, it is NP-hard. With this in mind, we use the one norm as a proxy for the zero norm and minimize $\|\mathbf{x}\|_1$ over $\{\mathbf{x} | A\mathbf{x} = \mathbf{b}\}$. Although this problem appears to be nonlinear, it can be solved by linear programming by writing $\mathbf{x} = \mathbf{u} - \mathbf{v}$, $\mathbf{u} \geq 0$, and $\mathbf{v} \geq 0$, and then minimizing the linear function $\sum_i u_i + \sum_i v_i$ subject to $A\mathbf{u} - A\mathbf{v} = \mathbf{b}$, $\mathbf{u} \geq 0$, and $\mathbf{v} \geq 0$.

Under what conditions will minimizing $\|\mathbf{x}\|_1$ over $\{\mathbf{x} | A\mathbf{x} = \mathbf{b}\}$ recover the s -sparse solution to $A\mathbf{x} = \mathbf{b}$? If $g(\mathbf{x})$ is a convex function, then any local minimum of g is a global minimum. If $g(\mathbf{x})$ is differentiable at its minimum, the gradient ∇g must be zero there. However, the 1-norm is not differentiable at its minimum. Thus, we introduce the concept of a subgradient of a convex function. Where the function is differentiable the subgradient is just the gradient. Where the function is not differentiable, the sub gradient is any line touching the function at the point that lies totally below the function. See Figure 10.3.

Subgradients are defined as follows: A *subgradient* of a function g at a point $\mathbf{x}_0$, is a vector $\nabla g(\mathbf{x}_0)$ satisfying $g(\mathbf{x}_0 + \Delta\mathbf{x}) \geq g(\mathbf{x}_0) + (\nabla g)^T \Delta\mathbf{x}$ for any vector $\Delta\mathbf{x}$. A point is a minimum for a convex function if there is a subgradient at that point with slope zero.

Consider the function $\|x\|_1$, where x is a real variable. For $x < 0$, the subgradient equals the gradient and has value -1. For $x > 0$, the subgradient equals the gradient and has value 1. At $x = 0$, the subgradient can be any value in the range $[-1, 1]$. The following proposition generalizes this example to the 1-norm function in d -space.

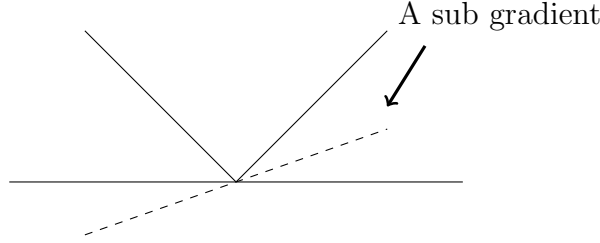


Figure 10.4: Illustration of a subgradient for $\|\mathbf{x}\|_1$ at $\mathbf{x} = 0$

Proposition 10.2 *A vector $\mathbf{v}$ is a subgradient of the 1-norm function $\|\mathbf{x}\|_1$ at $\mathbf{x}$ if and only if it satisfies the three conditions below:*

1. $v_i = -1$ for all i in I_1 where, $I_1 = \{i | x_i < 0\}$,
2. $v_i = 1$ for all i in I_2 where, $I_2 = \{i | x_i > 0\}$,
3. and v_i in $[-1, 1]$ for all i in I_3 where, $I_3 = \{i | x_i = 0\}$.

Proof: It is easy to see that for any vector $\mathbf{y}$,

$$\|\mathbf{x} + \mathbf{y}\|_1 - \|\mathbf{x}\|_1 \geq -\sum_{i \in I_1} y_i + \sum_{i \in I_2} y_i + \sum_{i \in I_3} |y_i|.$$

If $\mathbf{v}$ satisfies the conditions in the proposition, then $\|\mathbf{x} + \mathbf{y}\|_1 \geq \|\mathbf{x}\|_1 + \mathbf{v}^T \mathbf{y}$ as required. Now for the converse, suppose that $\mathbf{v}$ is a subgradient. Consider a vector $\mathbf{y}$ that is zero in all components except the first and y_1 is non-zero with $y_1 = \pm \varepsilon$ for a small $\varepsilon > 0$. If $1 \in I_1$, then, $\|\mathbf{x} + \mathbf{y}\|_1 - \|\mathbf{x}\|_1 = -y_1$ which implies that $-y_1 \geq v_1 y_1$. Choosing $y_1 = \varepsilon$, gives $-1 \geq v_1$ and choosing $y_1 = -\varepsilon$, gives $-1 \leq v_1$. So $v_1 = -1$. Similar reasoning gives the second condition. For the third condition, choose i in I_3 and set $y_i = \pm \varepsilon$ and argue similarly. ■

To characterize the value of $\mathbf{x}$ that minimizes $\|\mathbf{x}\|_1$ subject to $A\mathbf{x}=\mathbf{b}$, note that at the minimum $\mathbf{x}_0$, there can be no downhill direction consistent with the constraint $A\mathbf{x}=\mathbf{b}$. Thus, if the direction $\Delta\mathbf{x}$ at $\mathbf{x}_0$ is consistent with the constraint $A\mathbf{x}=\mathbf{b}$, that is $A\Delta\mathbf{x}=\mathbf{0}$ so that $A(\mathbf{x}_0 + \Delta\mathbf{x}) = \mathbf{b}$, any subgradient ∇ for $\|\mathbf{x}\|_1$ at $\mathbf{x}_0$ must satisfy $\nabla^T \Delta\mathbf{x} = 0$.

A sufficient but not necessary condition for $\mathbf{x}_0$ to be a minimum is that there exists some $\mathbf{w}$ such that the sub gradient at $\mathbf{x}_0$ is given by $\nabla = A^T \mathbf{w}$. Then for any $\Delta\mathbf{x}$ such that $A\Delta\mathbf{x} = \mathbf{0}$, we have $\nabla^T \Delta\mathbf{x} = \mathbf{w}^T A\Delta\mathbf{x} = \mathbf{w}^T \cdot \mathbf{0} = 0$. That is, for any direction consistent with the constraint $A\mathbf{x} = \mathbf{b}$, the subgradient is zero and hence $\mathbf{x}_0$ is a minimum.

10.3.2 The Exact Reconstruction Property

Theorem 10.3 below gives a condition that guarantees that a solution $\mathbf{x}_0$ to $A\mathbf{x} = \mathbf{b}$ is the unique minimum 1-norm solution to $A\mathbf{x} = \mathbf{b}$. This is a sufficient condition, but not necessary condition.

Theorem 10.3 Suppose $\mathbf{x}_0$ satisfies $A\mathbf{x}_0 = \mathbf{b}$. If there is a subgradient ∇ to the 1-norm function at $\mathbf{x}_0$ for which there exists a $\mathbf{w}$ where $\nabla = A^T \mathbf{w}$ and the columns of A corresponding to nonzero components of $\mathbf{x}_0$ are linearly independent, then, $\mathbf{x}_0$ minimizes $\|\mathbf{x}\|_1$ subject to $A\mathbf{x} = \mathbf{b}$. Furthermore, these conditions imply that $\mathbf{x}_0$ is the unique minimum.

Proof: We first show that $\mathbf{x}_0$ minimizes $\|\mathbf{x}\|_1$. Suppose $\mathbf{y}$ is another solution to $A\mathbf{x} = \mathbf{b}$. We need to show that $\|\mathbf{y}\|_1 \geq \|\mathbf{x}_0\|_1$. Let $\mathbf{z} = \mathbf{y} - \mathbf{x}_0$. Then $A\mathbf{z} = A\mathbf{y} - A\mathbf{x}_0 = \mathbf{0}$. Hence, $\nabla^T \mathbf{z} = (A^T \mathbf{w})^T \mathbf{z} = \mathbf{w}^T A\mathbf{z} = 0$. Now, since ∇ is a subgradient of the 1-norm function at $\mathbf{x}_0$,

$$\|\mathbf{y}\|_1 = \|\mathbf{x}_0 + \mathbf{z}\|_1 \geq \|\mathbf{x}_0\|_1 + \nabla^T \cdot \mathbf{z} = \|\mathbf{x}_0\|_1$$

and so we have that $\|\mathbf{x}_0\|_1$ minimizes $\|\mathbf{x}\|_1$ over all solutions to $A\mathbf{x} = \mathbf{b}$.

Suppose $\tilde{\mathbf{x}}_0$ were another minimum. Then ∇ is also a subgradient at $\tilde{\mathbf{x}}_0$ as it is at $\mathbf{x}_0$. To see this, for $\Delta\mathbf{x}$ such that $A\Delta\mathbf{x} = \mathbf{0}$,

$$\|\tilde{\mathbf{x}}_0 + \Delta\mathbf{x}\|_1 = \left\| \mathbf{x}_0 + \underbrace{\tilde{\mathbf{x}}_0 - \mathbf{x}_0 + \Delta\mathbf{x}}_{\alpha} \right\|_1 \geq \|\mathbf{x}_0\|_1 + \nabla^T (\tilde{\mathbf{x}}_0 - \mathbf{x}_0 + \Delta\mathbf{x}).$$

The above equation follows from the definition of ∇ being a subgradient for the one norm function, $\|\cdot\|_1$, at $\mathbf{x}_0$. Thus,

$$\|\tilde{\mathbf{x}}_0 + \Delta\mathbf{x}\|_1 \geq \|\mathbf{x}_0\|_1 + \nabla^T (\tilde{\mathbf{x}}_0 - \mathbf{x}_0) + \nabla^T \Delta\mathbf{x}.$$

But

$$\nabla^T (\tilde{\mathbf{x}}_0 - \mathbf{x}_0) = \mathbf{w}^T A (\tilde{\mathbf{x}}_0 - \mathbf{x}_0) = \mathbf{w}^T (\mathbf{b} - \mathbf{b}) = 0.$$

Hence, since $\tilde{\mathbf{x}}_0$ being a minimum means $\|\tilde{\mathbf{x}}_0\|_1 = \|\mathbf{x}_0\|_1$,

$$\|\tilde{\mathbf{x}}_0 + \Delta\mathbf{x}\|_1 \geq \|\mathbf{x}_0\|_1 + \nabla^T \Delta\mathbf{x} = \|\tilde{\mathbf{x}}_0\|_1 + \nabla^T \Delta\mathbf{x}.$$

This implies that ∇ is a sub gradient at $\tilde{\mathbf{x}}_0$.

Now, ∇ is a subgradient at both $\mathbf{x}_0$ and $\tilde{\mathbf{x}}_0$. By Proposition 10.2, we must have that $(\nabla)_i = \text{sgn}((x_0)_i) = \text{sgn}((\tilde{x}_0)_i)$, whenever either is nonzero and $|(\nabla)_i| < 1$, whenever either is 0. It follows that $\mathbf{x}_0$ and $\tilde{\mathbf{x}}_0$ have the same sparseness pattern. Since $A\mathbf{x}_0 = \mathbf{b}$ and $A\tilde{\mathbf{x}}_0 = \mathbf{b}$ and $\mathbf{x}_0$ and $\tilde{\mathbf{x}}_0$ are both nonzero on the same coordinates, and by the assumption that the columns of A corresponding to the nonzeros of $\mathbf{x}_0$ and $\tilde{\mathbf{x}}_0$ are independent, it must be that $\mathbf{x}_0 = \tilde{\mathbf{x}}_0$. ■

10.3.3 Restricted Isometry Property

Next we introduce the restricted isometry property that plays a key role in exact reconstruction of sparse vectors. A matrix A satisfies the *restricted isometry property*, *RIP*, if there exists a sequence of real numbers δ_s such that for any s -sparse $\mathbf{x}$

$$(1 - \delta_s) \|\mathbf{x}\|^2 \leq \|A\mathbf{x}\|^2 \leq (1 + \delta_s) \|\mathbf{x}\|^2. \quad (10.1)$$

Isometry is a mathematical concept; it refers to linear transformations that exactly preserve length such as rotations. If A is an $n \times n$ isometry, all its eigenvalues are ± 1 and it represents a coordinate system. Since a pair of orthogonal vectors are orthogonal in all coordinate system, for an isometry A and two orthogonal vectors $\mathbf{x}$ and $\mathbf{y}$, $\mathbf{x}^T A^T A \mathbf{y} = 0$. We will prove approximate versions of these properties for matrices A satisfying the restricted isometry property. The approximate versions will be used in the sequel.

A piece of notation will be useful. For a subset S of columns of A , let A_S denote the submatrix of A consisting of the columns of S .

Lemma 10.4 *If A satisfies the restricted isometry property, then*

1. *For any subset S of columns with $|S| = s$, the singular values of A_S are all between $1 - \delta_s$ and $1 + \delta_s$.*
2. *For any two orthogonal vectors $\mathbf{x}$ and $\mathbf{y}$, with supports of size s_1 and s_2 respectively, $|\mathbf{x}^T A^T A \mathbf{y}| \leq 5|\mathbf{x}||\mathbf{y}|(\delta_{s_1} + \delta_{s_2})$.*

Proof: Item 1 follows from the definition. To prove the second item, assume without loss of generality that $|\mathbf{x}| = |\mathbf{y}| = 1$. Since $\mathbf{x}$ and $\mathbf{y}$ are orthogonal, $|\mathbf{x} + \mathbf{y}|^2 = 2$. Consider $|A(\mathbf{x} + \mathbf{y})|^2$. This is between $2(1 - \delta_{s_1} + \delta_{s_2})^2$ and $2(1 + \delta_{s_1} + \delta_{s_2})^2$ by the restricted isometry property. Also $|A\mathbf{x}|^2$ is between $(1 - \delta_{s_1})^2$ and $(1 + \delta_{s_1})^2$ and $|A\mathbf{y}|^2$ is between $(1 - \delta_{s_2})^2$ and $(1 + \delta_{s_2})^2$. Since

$$\begin{aligned} 2\mathbf{x}^T A^T A \mathbf{y} &= (\mathbf{x} + \mathbf{y})^T A^T A (\mathbf{x} + \mathbf{y}) - \mathbf{x}^T A^T A \mathbf{x} - \mathbf{y}^T A^T A \mathbf{y} \\ &= |A(\mathbf{x} + \mathbf{y})|^2 - |A\mathbf{x}|^2 - |A\mathbf{y}|^2, \end{aligned}$$

it follows that

$$\begin{aligned} |2\mathbf{x}^T A^T A \mathbf{y}| &\leq 2(1 + \delta_{s_1} + \delta_{s_2})^2 - (1 - \delta_{s_1})^2 - (1 - \delta_{s_2})^2 \\ &6(\delta_{s_1} + \delta_{s_2}) + (\delta_{s_1}^2 + \delta_{s_2}^2 + 4\delta_{s_1} + 4\delta_{s_2}) \leq 9(\delta_{s_1} + \delta_{s_2}). \end{aligned}$$

Thus, for arbitrary $\mathbf{x}$ and $\mathbf{y}$ $|\mathbf{x}^T A^T A \mathbf{y}| \leq (9/2)|\mathbf{x}||\mathbf{y}|(\delta_{s_1} + \delta_{s_2})$. ■

Theorem 10.5 *Suppose A satisfies the restricted isometry property with*

$$\delta_{s+1} \leq \frac{1}{10\sqrt{s}}.$$

Suppose $\mathbf{x}_0$ has at most s nonzero coordinates and satisfies $A\mathbf{x} = \mathbf{b}$. Then, a subgradient $\nabla ||(\mathbf{x}_0)||_1$ for the 1-norm function exists at $\mathbf{x}_0$ which satisfies the conditions of Theorem 10.3 and so $\mathbf{x}_0$ is the unique minimum 1-norm solution to $A\mathbf{x} = \mathbf{b}$.

Proof: Let

$$S = \{i | (\mathbf{x}_0)_i \neq 0\}$$

be the support of $\mathbf{x}_0$ and let $\bar{S} = \{i | (\mathbf{x}_0)_i = 0\}$ be the complement set of coordinates. To find a subgradient $\mathbf{u}$ at $\mathbf{x}_0$ satisfying Theorem 10.3, search for a $\mathbf{w}$ such that $\mathbf{u} = A^T \mathbf{w}$ where for coordinates in which $\mathbf{x}_0 \neq 0$, $\mathbf{u} = \text{sgn}(\mathbf{x}_0)$ and for the remaining coordinates the 2-norm of $\mathbf{u}$ is minimized. Solving for $\mathbf{w}$ is a least squares problem. Let $\mathbf{z}$ be the vector with support S , with $z_i = \text{sgn}(\mathbf{x}_0)$ on S . Consider the vector $\mathbf{w}$ defined by

$$\mathbf{w} = A_S (A_S^T A_S)^{-1} \mathbf{z}.$$

This happens to be the solution of the least squares problem, but we do not need this fact. We only state it to tell the reader how we came up with this expression. Note that A_S has independent columns from the restricted isometry property assumption, and so $A_S^T A_S$ is invertible. We will prove that this $\mathbf{w}$ satisfies the conditions of Theorem 10.3. First, for coordinates in S , we have

$$(A^T \mathbf{w})_S = (A_S)^T A_S (A_S^T A_S)^{-1} \mathbf{z} = \mathbf{z}$$

as required.

For coordinates in $\bar{S}$, we have

$$(A^T \mathbf{w})_{\bar{S}} = (A_{\bar{S}})^T A_S (A_S^T A_S)^{-1} \mathbf{z}.$$

Now, the eigenvalues of $A_S^T A_S$, which are the squares of the singular values of A_S , are between $(1 - \delta_s)^2$ and $(1 + \delta_s)^2$. So $\|(A_S^T A_S)^{-1}\| \leq \frac{1}{(1 - \delta_s)^2}$. Letting $\mathbf{p} = (A_S^T A_S)^{-1} \mathbf{z}$, we have $|\mathbf{p}| \leq \frac{\sqrt{s}}{(1 - \delta_s)^2}$. Write $A_s \mathbf{p}$ as $A \mathbf{q}$, where $\mathbf{q}$ has all coordinates in $\bar{S}$ equal to zero. Now, for $j \in \bar{S}$

$$(A^T \mathbf{w})_j = e_j^T A^T A \mathbf{q}$$

and part (2) of Lemma 10.4 gives $|(A^T \mathbf{w})_j| \leq 9\delta_{s+1}\sqrt{s}/(1 - \delta_s^2) \leq 1/2$ establishing the Theorem 10.3 holds. ■

A Gaussian matrix is a matrix where each element is an independent gaussian variable. Gaussian matrices satisfy the restricted isometry property. (Exercise ??)

10.4 Applications

10.4.1 Sparse Vector in Some Coordinate Basis

Consider $A\mathbf{x} = \mathbf{b}$ where A is a square $n \times n$ matrix. The vectors $\mathbf{x}$ and $\mathbf{b}$ can be considered as two representations of the same quantity. For example, $\mathbf{x}$ might be a discrete time sequence with $\mathbf{b}$ the frequency spectrum of $\mathbf{x}$ and the matrix A the Fourier transform. The quantity $\mathbf{x}$ can be represented in the time domain by $\mathbf{x}$ and in the frequency domain by its Fourier transform $\mathbf{b}$. In fact, any orthonormal matrix can be thought of as a transformation and there are many important transformations other than the Fourier transformation.

Consider a transformation A and a signal $\mathbf{x}$ in some standard representation. Then $\mathbf{y} = A\mathbf{x}$ transforms the signal $\mathbf{x}$ to another representation $\mathbf{y}$. If A spreads any sparse signal $\mathbf{x}$ out so that the information contained in each coordinate in the standard basis is spread out to all coordinates in the second basis, then the two representations are said to be *incoherent*. A signal and its Fourier transform are one example of incoherent vectors. This suggests that if $\mathbf{x}$ is sparse, only a few randomly selected coordinates of its Fourier transform are needed to reconstruct $\mathbf{x}$. In the next section we show that a signal cannot be too sparse in both its time domain and its frequency domain.

10.4.2 A Representation Cannot be Sparse in Both Time and Frequency Domains

We now show that there is an uncertainty principle that states that a time signal cannot be sparse in both the time domain and the frequency domain. If the signal is of length n , then the product of the number of nonzero coordinates in the time domain and the number of nonzero coordinates in the frequency domain must be at least n . We first prove two technical lemmas.

In dealing with the Fourier transform it is convenient for indices to run from 0 to $n-1$ rather than from 1 to n . Let $x_0, x_1, \dots, x_{n-1}$ be a sequence. $f_0, f_1, \dots, f_{n-1}$ be its discrete Fourier transform. Let $i = \sqrt{-1}$. Then $f_j = \frac{1}{\sqrt{n}} \sum_{k=0}^{n-1} x_k e^{-\frac{2\pi i}{n} jk}$, $j = 0, \dots, n-1$. In matrix form $\mathbf{f} = Z\mathbf{x}$ where $z_{jk} = e^{-\frac{2\pi i}{n} jk}$.

$$\begin{pmatrix} f_0 \\ f_1 \\ \vdots \\ f_{n-1} \end{pmatrix} = \frac{1}{\sqrt{n}} \begin{pmatrix} 1 & 1 & \dots & 1 \\ e^{-\frac{2\pi i}{n}} & e^{-\frac{2\pi i}{n} 2} & \dots & e^{-\frac{2\pi i}{n} (n-1)} \\ \vdots & \vdots & & \vdots \\ e^{-\frac{2\pi i}{n} (n-1)} & e^{-\frac{2\pi i}{n} 2(n-1)} & \dots & e^{-\frac{2\pi i}{n} (n-1)^2} \end{pmatrix} \begin{pmatrix} x_0 \\ x_1 \\ \vdots \\ x_{n-1} \end{pmatrix}$$

If some of the elements of $\mathbf{x}$ are zero, delete the corresponding columns from the matrix. To maintain a square matrix, let n_x be the number of nonzero elements in $\mathbf{x}$ and select n_x consecutive rows of Z . Normalize the columns of the resulting submatrix by dividing each element in a column by the column element in the first row. The resulting submatrix is a Vandermonde matrix that looks like

$$\begin{pmatrix} 1 & 1 & 1 & 1 \\ a & b & c & d \\ a^2 & b^2 & c^2 & d^2 \\ a^3 & b^3 & c^3 & d^3 \end{pmatrix}$$

which is nonsingular.

Lemma 10.6 *If $x_0, x_1, \dots, x_{n-1}$ has n_x nonzero elements, then $f_0, f_1, \dots, f_{n-1}$ cannot have n_x consecutive zeros.*

Proof: Let $i_1, i_2, \dots, i_{n_x}$ be the indices of the nonzero elements of $\mathbf{x}$. Then the elements of the Fourier transform in the range $k = m + 1, m + 2, \dots, m + n_x$ are

$$f_k = \frac{1}{\sqrt{n}} \sum_{j=1}^{n_x} x_{i_j} e^{\frac{-2\pi i}{n} k i_j}$$

Note the use of i as $\sqrt{-1}$ and the multiplication of the exponent by i_j to account for the actual location of the element in the sequence. Normally, if every element in the sequence was included, we would just multiply by the index of summation.

Convert the equation to matrix form by defining $z_{kj} = \frac{1}{\sqrt{n}} \exp(-\frac{2\pi i}{n} k i_j)$ and write $\mathbf{f} = Z\mathbf{x}$. Actually instead of $\mathbf{x}$, write the vector consisting of the nonzero elements of $\mathbf{x}$. By its definition, $\mathbf{x} \neq 0$. To prove the lemma we need to show that $\mathbf{f}$ is nonzero. This will be true provided Z is nonsingular. If we rescale Z by dividing each column by its leading entry we get the Vandermonde determinant which is nonsingular. ■

Theorem 10.7 *Let n_x be the number of nonzero elements in $\mathbf{x}$ and let n_f be the number of nonzero elements in the Fourier transform of $\mathbf{x}$. Let n_x divide n . Then $n_x n_f \geq n$.*

Proof: If $\mathbf{x}$ has n_x nonzero elements, $\mathbf{f}$ cannot have a consecutive block of n_x zeros. Since n_x divides n there are $\frac{n}{n_x}$ blocks each containing at least one nonzero element. Thus, the product of nonzero elements in $\mathbf{x}$ and $\mathbf{f}$ is at least n . ■

Fourier transform of spikes prove that above bound is tight

To show that the bound in Theorem 10.7 is tight we show that the Fourier transform of the sequence of length n consisting of $\sqrt{n}$ ones, each one separated by $\sqrt{n} - 1$ zeros, is the sequence itself. For example, the Fourier transform of the sequence 100100100 is 100100100. Thus, for this class of sequences, $n_x n_f = n$. The Fourier transform is

$$y_j = \frac{1}{\sqrt{n}} \sum_{k=0}^{n-1} z^{jk}.$$

Theorem 10.8 *Let $S(\sqrt{n}, \sqrt{n})$ be the sequence of 1's and 0's with $\sqrt{n}$ 1's spaced $\sqrt{n}$ apart. The Fourier transform of $S(\sqrt{n}, \sqrt{n})$ is itself.*

Proof: Consider the columns $0, \sqrt{n}, 2\sqrt{n}, \dots, (\sqrt{n} - 1)\sqrt{n}$. These are the columns for which $S(\sqrt{n}, \sqrt{n})$ has value 1. The element of the matrix Z in the row $j\sqrt{n}$ of column $k\sqrt{n}$, $0 \leq k < \sqrt{n}$ is $z^{nkj} = 1$. Thus, for these rows Z times the vector $S(\sqrt{n}, \sqrt{n}) = \sqrt{n}$ and the $1/\sqrt{n}$ normalization yields $f_{j\sqrt{n}} = 1$.

For rows whose index is not of the form $j\sqrt{n}$, the row b , $b \neq j\sqrt{n}$, $j \in \{0, \sqrt{n}, \dots, \sqrt{n} - 1\}$, the elements in row b in the columns $0, \sqrt{n}, 2\sqrt{n}, \dots, (\sqrt{n} - 1)\sqrt{n}$ are $1, z^b, z^{2b}, \dots, z^{(\sqrt{n}-1)b}$ and thus $f_b = \frac{1}{\sqrt{n}} \left(1 + z^b + z^{2b} \dots + z^{(\sqrt{n}-1)b} \right) = \frac{1}{\sqrt{n}} \frac{z^{\sqrt{n}b} - 1}{z^b - 1} = 0$ since $z^{b\sqrt{n}} = 1$ and $z \neq 1$.

1	1	1	1	1	1	1	1	1
1	z	z^2	z^3	z^4	z^5	z^6	z^7	z^8
1	z^2	z^4	z^6	z^8	z	z^3	z^5	z^7
1	z^3	z^6	1	z^3	z^6	1	z^3	z^6
1	z^4	z^8	z^3	z^7	z^2	z^6	z	z^5
1	z^5	z	z^6	z^2	z^7	z^3	z^8	z^4
1	z^6	z^3	1	z^6	z^3	1	z^6	z^3
1	z^7	z^5	z^3	z	z^8	z^6	z^4	z^2
1	z^8	z^7	z^6	z^5	z^4	z^3	z^2	z

Figure 10.5: The matrix Z for $n=9$.

Uniqueness of l_1 optimization

Consider a redundant representation for a sequence. One such representation would be representing a sequence as the concatenation of two sequences, one specified by its coordinates and the other by its Fourier transform. Suppose some sequence could be represented as a sequence of coordinates and Fourier coefficients sparsely in two different ways. Then by subtraction, the zero sequence could be represented by a sparse sequence. The representation of the zero sequence cannot be solely coordinates or Fourier coefficients. If y is the coordinate sequence in the representation of the zero sequence, then the Fourier portion of the representation must represent $-y$. Thus y and its Fourier transform would have sparse representations contradicting $n_x n_f \geq n$. Notice that a factor of two comes in when we subtract the two representations.

Suppose two sparse signals had Fourier transforms that agreed in almost all of their coordinates. Then the difference would be a sparse signal with a sparse transform. This is not possible. Thus, if one selects $\log n$ elements of their transform these elements should distinguish between these two signals.

10.4.3 Biological

There are many areas where linear systems arise in which a sparse solution is unique. One is in plant breeding. Consider a breeder who has a number of apple trees and for each tree observes the strength of some desirable feature. He wishes to determine which genes are responsible for the feature so he can cross breed to obtain a tree that better expresses the desirable feature. This gives rise to a set of equations $A\mathbf{x} = \mathbf{b}$ where each row of the matrix A corresponds to a tree and each column to a position on the genome. See Figure 10.6. The vector $\mathbf{b}$ corresponds to the strength of the desired feature in each tree. The solution $\mathbf{x}$ tells us the position on the genome corresponding to the genes that account for the feature. It would be surprising if there were two small independent sets of genes that accounted for the desired feature. Thus, the matrix must have a property that allows only one sparse solution.

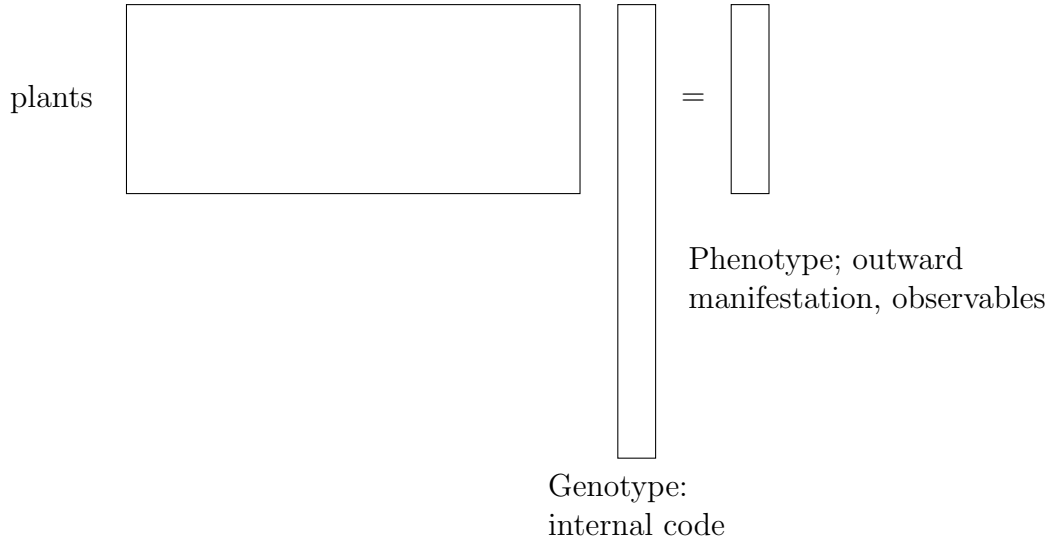


Figure 10.6: The system of linear equations used to find the internal code for some observable phenomenon.

10.4.4 Finding Overlapping Cliques or Communities

Consider a graph that consists of several cliques. Suppose we can observe only low level information such as edges and we wish to identify the cliques. An instance of this problem is the task of identifying which of ten players belongs to which of two teams of five players each when one can only observe interactions between pairs of individuals. There is an interaction between two players if and only if they are on the same team. In this situation we have a matrix A with $\binom{10}{5}$ columns and $\binom{10}{2}$ rows. The columns represent possible teams and the rows represent pairs of individuals. Let $\mathbf{b}$ be the $\binom{10}{2}$ dimensional vector of observed interactions. Let $\mathbf{x}$ be a solution to $A\mathbf{x} = \mathbf{b}$. There is a sparse solution $\mathbf{x}$ where $\mathbf{x}$ is all zeros except for the two 1's for 12345 and 678910 where the two teams are $\{1,2,3,4,5\}$ and $\{6,7,8,9,10\}$. The question is can we recover $\mathbf{x}$ from $\mathbf{b}$. If the matrix A had satisfied the restricted isometry condition, then we could surely do this. Although A does not satisfy the restricted isometry condition which guarantees recover of all sparse vectors, we can recover the sparse vector in the case where the teams are non overlapping or almost non overlapping. If A satisfied the restricted isometry property we would minimize $\|\mathbf{x}\|_1$ subject to $A\mathbf{x} = \mathbf{b}$. Instead, we minimize $\|\mathbf{x}\|_1$ subject to $\|A\mathbf{x} - \mathbf{b}\|_\infty \leq \varepsilon$ where we bound the largest error.

10.4.5 Low Rank Matrices

Suppose we have a low rank matrix L that has been corrupted by noise. That is $M = L + R$. If the R is Gaussian, then principle component analysis will recover L from M . However, if L has been corrupted by several missing entries or several entries have a

large noise added to them and they become outliers, then principle component analysis may be far off. However, if L is low rank and R is sparse, then we can recover L effectively from $L + R$. To do this we find the L and R that minimize $\|L\|_* + \lambda \|R\|_1$. Here $\|L\|_*$ is the sum of the singular values of L . Notice that we do not need to know the rank of L or the elements that were corrupted. All we need is that the low rank matrix L is not sparse and that the sparse matrix R is not low rank. We leave the proof as an exercise.

An example where low rank matrices that have been corrupted might occur is aerial photographs of an intersection. Given a long sequence of such photographs, they will be the same except for cars and people. If each photo is converted to a vector and the vector used to make a column of a matrix, then the matrix will be low rank corrupted by the traffic. Finding the original low rank matrix will separate the cars and people from the back ground.

10.5 Exercises

Exercise 10.1 *Select a method of combining individual rankings into a global ranking. Consider a set of rankings where each individual ranks b last. One by one move b from the bottom to the top leaving the other rankings in place. Determine v_b as in Theorem 10.1 where b_b is the ranking that causes b to move from the bottom to the top in the global ranking.*

Exercise 10.2 *Show that the three axioms: non dictator, unanimity, and independence of irrelevant alternatives are independent.*

Exercise 10.3 *Does the axiom of independence of irrelevant alternatives make sense? What if there were three rankings of five items. In the first two rankings, A is number one and B is number two. In the third ranking, B is number one and A is number five. One might compute an average score where a low score is good. A gets a score of $1+1+5=7$ and B gets a score of $2+2+1=5$ and B is ranked number one in the global ranking. Now if the third raker moves A up to the second position, A 's score becomes $1+1+2=4$ and the global ranking of A and B changes even though no individual ranking of A and B changed. Is there some alternative axiom to replace independence of irrelevant alternatives?*

Exercise 10.4 *Prove that the global ranking agrees with column v_b even if b is moved down through the column.*

Exercise 10.5 *Create a random 100 by 100 orthonormal matrix A and a sparse 100-dimensional vector $\mathbf{b}$. Compute $\mathbf{y} = A\mathbf{x}$. Randomly select a few coordinates of $\mathbf{y}$ and reconstruct $\mathbf{x}$ from the samples of $\mathbf{y}$ using the minimization of 1-norm technique of Section 10.3.1. Did you get $\mathbf{x}$ back?*

Exercise 10.6 (maybe belongs in a different chapter) *Let A be a low rank $n \times m$ matrix. Let r be the rank of A . Let $\tilde{A}$ be A corrupted by Gaussian noise. Prove that the rank r SVD approximation to $\tilde{A}$ minimizes $\|A - \tilde{A}\|_F^2$.*

Exercise 10.7 *Prove that minimizing $\|x\|_0$ subject to $Ax = b$ is NP-complete.*

Exercise 10.8 *Let A be a Gaussian matrix where each element is a random Gaussian variable with zero mean and variance one. Prove that A has the restricted isometry property.*

Exercise 10.9 *Generate 100×100 matrices of rank 20, 40, 60, 80, and 100. In each matrix randomly delete 50, 100, 200, or 400 entries. In each case try to recover the original matrix. How well do you do?*

Exercise 10.10 *Repeat the previous exercise but instead of deleting elements, corrupt the elements by adding a reasonable size corruption to the randomly selected matrix entries.*

Exercise 10.11 Compute the Fourier transform of the sequence 1000010000.

Exercise 10.12 What is the Fourier transform of a cyclic shift?

Exercise 10.13 The number $n=6$ is factorable but not a perfect square. What is Fourier transform of $F(2,3)$? Is the transform of 100100, 101010?

Exercise 10.14 Let Z be the n root of unity. Prove that $\{z^{bi} | 0 \leq i < n\} = \{z^i | 0 \leq i < n\}$ proved that b does not divide n .

Exercise 10.15 The Vandermonde determinant is of the form

$$\begin{pmatrix} 1 & 1 & 1 & 1 \\ a & b & c & d \\ a^2 & b^2 & c^2 & d^2 \\ a^3 & b^3 & c^3 & d^3 \end{pmatrix}$$

Show that if a, b, c , and d are distinct, then the Vandermonde determinant is nonsingular.

Hint: Given value at each of n points, there is a unique polynomial that takes on the values at the points.

Exercise 10.16 Many problems can be formulated as finding $\mathbf{x}$ satisfying $A\mathbf{x} = \mathbf{b} + \mathbf{r}$ where $\mathbf{r}$ is some residual error. Discuss the advantages and disadvantages of each of the following three versions of the problem.

1. Set $\mathbf{r}=0$ and find $\mathbf{x} = \operatorname{argmin} \|\mathbf{x}\|_1$ satisfying $A\mathbf{x} = \mathbf{b}$
2. Lasso: find $\mathbf{x} = \operatorname{argmin} (\|\mathbf{x}\|_1 + \alpha \|\mathbf{r}\|_2^2)$ satisfying $A\mathbf{x} = \mathbf{b}$
3. find $\underline{x} = \operatorname{argmin} \|\mathbf{x}\|_1$ such that $\|\mathbf{r}\|_2 < \varepsilon$

Exercise 10.17 Create a graph of overlapping communities as follows. Let $n=1,000$. Partition the integers into ten blocks each of size 100. The first block is $\{1, 2, \dots, 100\}$. The second is $\{100, 101, \dots, 200\}$ and so on. Add edges to the graph so that the vertices in each block form a clique. Now randomly permute the indices and partition the sequence into ten blocks of 100 vertices each. Again add edges so that these new blocks are cliques. Randomly permute the indices a second time and repeat the process of adding edges. The result is a graph in which each vertex is in three cliques. Explain how to find the cliques given the graph.

Exercise 10.18 Repeat the above exercise but instead of adding edges to form cliques, use each block to form a $G(100, p)$ graph. For how small a p can you recover the blocks? What if you add $G(1,000, q)$ to the graph for some small value of q .

Exercise 10.19 Construct an $n \times m$ matrix A where each of the m columns is a 0-1 indicator vector with approximately $1/4$ entries being 1. Then $B = AA^T$ is a symmetric matrix that can be viewed as the adjacency matrix of an n vertex graph. Some edges will have weight greater than one. The graph consists of a number of possibly overlapping cliques. Your task given B is to find the cliques by the following technique of finding a 0-1 vector in the column space of B by the following linear program for finding b and x .

$$b = \operatorname{argmin} \|b\|_1$$

subject to

$$Bx = b$$

$$b_1 = 1$$

$$0 \leq b_i \leq 1 \quad 2 \leq i \leq n$$

Then subtract bb^T from B and repeat.

Exercise 10.20 Construct an example of a matrix A satisfying the following conditions

1. The columns of A are 0-1 vectors where the support of no two columns overlap by 50% or more.
2. No column's support is totally within the support of another column.
3. The minimum 1-norm vector in the column space of A is not a 0-1 vector.

Exercise 10.21 Let $M = L + R$ where L is a low rank matrix corrupted by a sparse noise matrix R . Why can we not recover L from M if R is low rank or if L is sparse?

Exercise 10.22 LASSO Create exercise on LASSO $\min (\|Ax - b\|^2 + \lambda \|x\|_1)$

References

Arrow's impossibility theorem – Wikipedia

John Geanakoplos, "Three brief proofs of Arrow's Impossibility Theorem"

see <http://bcn.boulder.co.us/government/approvalvote/altvote.html> for Hare system

1. Xiaoye Jiang, Yuan Yao, and Leonidas J. Guibas, "Stable Identification of Cliques With Restricted Sensing"

2. Emmanuel J. Candes, "Compressive Sampling"

3. David L. Donoho and Philip B. Stark, "Uncertainty Principles and Signal Recovery"

4. David L. Donoho and Xiaoming Huo, "Uncertainty Principles and Ideal Atomic Decomposition"

5. Emmanuel J. Candes and Michael B. Wakin, "An Introduction to Compressive Sampling"

6. Emmanuel J. Candes, Xiaodong Li, Yi Ma, and John Wright, "Robust Principal Component Analysis"

11 Appendix

11.1 Asymptotic Notation

We introduce the big O notation here. The motivating example is the analysis of the running time of an algorithm. The running time may be a complicated function of the input length n such as $5n^3 + 25n^2 \ln n - 6n + 22$. Asymptotic analysis is concerned with the behavior as $n \rightarrow \infty$ where the higher order term $5n^3$ dominates. Further, the coefficient 5 of $5n^3$ is not of interest since its value varies depending on the machine model. So we say that the function is $O(n^3)$. The big O notation applies to functions on the positive integers taking on positive real values.

Definition 11.1 For functions f and g from the natural numbers to the positive reals, $f(n)$ is $O(g(n))$ if there exists a constant $c > 0$ such that for all n , $f(n) \leq cg(n)$. ■

Thus, $f(n) = 5n^3 + 25n^2 \ln n - 6n + 22$ is $O(n^3)$. The upper bound need not be tight. Not only is $f(n)$, $O(n^3)$, it is also $O(n^4)$. Note $g(n)$ must be strictly greater than 0 for all n .

To say that the function $f(n)$ grows at least as fast as $g(n)$, one uses a notation called omega of n . For positive real valued f and g , $f(n)$ is $\Omega(g(n))$ if there exists a constant $c > 0$ such that for all n , $f(n) \geq cg(n)$. If $f(n)$ is both $O(g(n))$ and $\Omega(g(n))$, then $f(n)$ is $\Theta(g(n))$. Theta of n is used when the two functions have the same asymptotic growth rate.

Many times one wishes to bound the low order terms. To do this, a notation called little o of n is used. We say $f(n)$ is $o(g(n))$ if $\lim_{n \rightarrow \infty} \frac{f(n)}{g(n)} = 0$. Note that $f(n)$ being $O(g(n))$ means that asymptotically $f(n)$ does not grow faster than $g(n)$, whereas $f(n)$ being $o(g(n))$ means that asymptotically $f(n)/g(n)$ goes to zero. If $f(n) = 2n + \sqrt{n}$, then

asymptotic upper bound	
$f(n)$ is $O(g(n))$ if for all n , $f(n) \leq cg(n)$ for some constant $c > 0$.	$\leq$
asymptotic lower bound	
$f(n)$ is $\Omega(g(n))$ if for all n , $f(n) \geq cg(n)$ for some constant $c > 0$.	$\geq$
asymptotic equality	
$f(n)$ is $\Theta(g(n))$ if it is both $O(g(n))$ and $\Omega(g(n))$.	$=$
$f(n)$ is $o(g(n))$ if $\lim_{n \rightarrow \infty} \frac{f(n)}{g(n)} = 0$.	$<$
$f(n) \sim g(n)$ if $\lim_{n \rightarrow \infty} \frac{f(n)}{g(n)} = 1$.	$=$
$f(n)$ is $\omega(g(n))$ if $\lim_{n \rightarrow \infty} \frac{f(n)}{g(n)} = \infty$.	$>$

$f(n)$ is $O(n)$ but in bounding the lower order term, we write $f(n) = 2n + o(n)$. Finally we write $f(n) \sim g(n)$ if $\lim_{n \rightarrow \infty} \frac{f(n)}{g(n)} = 1$ and say $f(n)$ is $\omega(g(n))$ if $\lim_{n \rightarrow \infty} \frac{f(n)}{g(n)} = \infty$. The difference between $f(n)$ being $\Theta(g(n))$ and $f(n) \sim g(n)$ is that in the first case $f(n)$ and $g(n)$ may differ by a multiplicative constant factor.

11.2 Sums of Series

Summations

$$\begin{aligned}\sum_{i=0}^n a^i &= 1 + a + a^2 + \cdots = \frac{1 - a^{n+1}}{1 - a}, \quad a \neq 1 \\ \sum_{i=0}^{\infty} a^i &= 1 + a + a^2 + \cdots = \frac{1}{1 - a}, \quad |a| < 1 \\ \sum_{i=0}^{\infty} i a^i &= a + 2a^2 + 3a^3 \cdots = \frac{a}{(1 - a)^2}, \quad |a| < 1 \\ \sum_{i=0}^{\infty} i^2 a^i &= a + 4a^2 + 9a^3 \cdots = \frac{a(1 + a)}{(1 - a)^3}, \quad |a| < 1 \\ \sum_{i=1}^n i &= \frac{n(n + 1)}{2} \\ \sum_{i=1}^n i^2 &= \frac{n(n + 1)(2n + 1)}{6} \\ \sum_{i=1}^{\infty} \frac{1}{i^2} &= \frac{\pi^2}{6}\end{aligned}$$

We prove one equality.

$$\sum_{i=0}^{\infty} i a^i = a + 2a^2 + 3a^3 \cdots = \frac{a}{(1 - a)^2}, \text{ provided } |a| < 1.$$

Write $S = \sum_{i=0}^{\infty} i a^i$.

$$aS = \sum_{i=0}^{\infty} i a^{i+1} = \sum_{i=1}^{\infty} (i - 1) a^i.$$

Thus,

$$S - aS = \sum_{i=1}^{\infty} i a^i - \sum_{i=1}^{\infty} (i - 1) a^i = \sum_{i=1}^{\infty} a^i = \frac{a}{1 - a},$$

from which the equality follows. The sum $\sum_i i^2 a^i$ can also be done by an extension of this method (left to the reader). Using generating functions, we will see another proof of both

these equalities by derivatives.

$$\sum_{i=1}^{\infty} \frac{1}{i} = 1 + \frac{1}{2} + \left(\frac{1}{3} + \frac{1}{4}\right) + \left(\frac{1}{5} + \frac{1}{6} + \frac{1}{7} + \frac{1}{8}\right) + \cdots \geq 1 + \frac{1}{2} + \frac{1}{2} + \cdots \text{ and thus diverges.}$$

The summation $\sum_{i=1}^n \frac{1}{i}$ grows as $\ln n$ since $\sum_{i=1}^n \frac{1}{i} \approx \int_{x=1}^n \frac{1}{x} dx$. In fact, $\lim_{i \rightarrow \infty} \left(\sum_{i=1}^n \frac{1}{i} - \ln(n) \right) = \gamma$ where $\gamma \cong 0.5772$ is Euler's constant. Thus, $\sum_{i=1}^n \frac{1}{i} \cong \ln(n) + \gamma$ for large n .

Truncated Maclaurin series

If all the derivatives of a function $f(x)$ exist, then we can write

$$f(x) = f(0) + f'(0)x + f''(0)\frac{x^2}{2} + \cdots.$$

The series can be truncated. In fact, there exists some y between 0 and x such that

$$f(x) = f(0) + f'(y)x.$$

Also, there exists some z between 0 and x such that

$$f(x) = f(0) + f'(0)x + f''(z)\frac{x^2}{2}$$

and so on for higher derivatives. This can be used to derive inequalities. For example, if $f(x) = \ln(1+x)$, then its derivatives are

$$f'(x) = \frac{1}{1+x} ; f''(x) = -\frac{1}{(1+x)^2} ; f'''(x) = \frac{2}{(1+x)^3}.$$

Thus, for any z , $f''(z) < 0$ and so we have for any x , $\ln(1+x) \leq x$ (which also follows from the inequality $1+x \leq e^x$). Also for $z > -1$, $f'''(z) > 0$ and so we get for $x > -1$,

$$\ln(1+x) > x - \frac{x^2}{2}.$$

Exponentials and logs

$$a^{\log n} = n^{\log a}$$

$$e^x = 1 + x + \frac{x^2}{2!} + \frac{x^3}{3!} + \cdots \quad e = 2.7182 \quad \frac{1}{e} = 0.3679$$

Setting $x = 1$ in $e^x = 1 + x + \frac{x^2}{2!} + \frac{x^3}{3!} + \cdots$ yields $e = \sum_{i=0}^{\infty} \frac{1}{i!}$.

$$\lim_{n \rightarrow \infty} \left(1 + \frac{a}{n}\right)^n = e^a$$

$$\ln(1+x) = x - \frac{1}{2}x^2 + \frac{1}{3}x^3 - \frac{1}{4}x^4 \cdots \quad |x| < 1$$

The above expression with $-x$ substituted for x gives rise to the approximations

$$\ln(1-x) < -x$$

which also follows from $1-x \leq e^{-x}$, since $\ln(1-x)$ is a monotone function for $x \in (0, 1)$.

For $0 < x < 0.69$, $\ln(1-x) > -x - x^2$.

Trigonometric identities

$$\sin(x \pm y) = \sin(x) \cos(y) \pm \cos(x) \sin(y)$$

$$\cos(x \pm y) = \cos(x) \cos(y) \mp \sin(x) \sin(y)$$

$$\cos(2\theta) = \cos^2 \theta - \sin^2 \theta = 1 - 2\sin^2 \theta$$

$$\sin(2\theta) = 2\sin \theta \cos \theta$$

$$\sin^2 \frac{\theta}{2} = \frac{1}{2}(1 - \cos \theta)$$

$$\cos^2 \frac{\theta}{2} = \frac{1}{2}(1 + \cos \theta)$$

Gaussian and related integrals

$$\int x e^{ax^2} dx = \frac{1}{2a} e^{ax^2}$$

$$\int \frac{1}{a^2+x^2} dx = \frac{1}{a} \tan^{-1} \frac{x}{a} \text{ thus } \int_{-\infty}^{\infty} \frac{1}{a^2+x^2} dx = \frac{\pi}{a}$$

$$\int_{-\infty}^{\infty} e^{-\frac{a^2 x^2}{2}} dx = \frac{\sqrt{2\pi}}{a} \text{ thus } \frac{a}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-\frac{a^2 x^2}{2}} dx = 1$$

$$\int_0^{\infty} x^2 e^{-ax^2} dx = \frac{1}{4a} \sqrt{\frac{\pi}{a}}$$

$$\int_0^{\infty} x^{2n} e^{-\frac{x^2}{a^2}} dx = \sqrt{\pi} \frac{1 \cdot 3 \cdot 5 \cdots (2n-1)}{2^{n+1}} a^{2n-1} = \sqrt{\pi} \frac{(2n)!}{n!} \left(\frac{a}{2}\right)^{2n+1}$$

$$\int_0^{\infty} x^{2n+1} e^{-\frac{x^2}{a^2}} dx = \frac{n!}{2} a^{2n+2}$$

$$\int_{-\infty}^{\infty} e^{-x^2} dx = \sqrt{\pi}$$

To verify $\int_{-\infty}^{\infty} e^{-x^2} dx = \sqrt{\pi}$, consider $\left(\int_{-\infty}^{\infty} e^{-x^2} dx \right)^2 = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-(x^2+y^2)} dx dy$. Let $x = r \cos \theta$ and $y = r \sin \theta$. The Jacobian of this transformation of variables is

$$J(r, \theta) = \left| \begin{array}{cc} \frac{\partial x}{\partial r} & \frac{\partial x}{\partial \theta} \\ \frac{\partial y}{\partial r} & \frac{\partial y}{\partial \theta} \end{array} \right| = \left| \begin{array}{cc} \cos \theta & -r \sin \theta \\ \sin \theta & r \cos \theta \end{array} \right| = r$$

Thus,

$$\begin{aligned} \left(\int_{-\infty}^{\infty} e^{-x^2} dx \right)^2 &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-(x^2+y^2)} dx dy = \int_0^{\infty} \int_0^{2\pi} e^{-r^2} J(r, \theta) dr d\theta \\ &= \int_0^{\infty} e^{-r^2} r dr \int_0^{2\pi} d\theta \\ &= -2\pi \left[\frac{e^{-r^2}}{2} \right]_0^{\infty} = \pi \end{aligned}$$

Thus, $\int_{-\infty}^{\infty} e^{-x^2} dx = \sqrt{\pi}$.

Miscellaneous integrals

$$\int_{x=0}^1 x^{\alpha-1} (1-x)^{\beta-1} dx = \frac{\Gamma(\alpha)\Gamma(\beta)}{\Gamma(\alpha+\beta)}$$

Binomial coefficients

The binomial coefficient $\binom{n}{k} = \frac{n!}{(n-k)!k!}$ is the number of ways of choosing k items from n .

$$\binom{n}{d} + \binom{n}{d+1} = \binom{n+1}{d+1}.$$

The observation that the number of ways of choosing k items from $2n$ equals the number of ways of choosing i items from the first n and choosing $k-i$ items from the second n summed over all i , $0 \leq i \leq k$ yields the identity

$$\sum_{i=0}^k \binom{n}{i} \binom{n}{k-i} = \binom{2n}{k}.$$

Setting $k = n$ in the above formula and observing that $\binom{n}{i} = \binom{n}{n-i}$ yields

$$\sum_{i=0}^n \binom{n}{i}^2 = \binom{2n}{n}$$

More generally $\sum_{i=0}^k \binom{n}{i} \binom{m}{k-i} = \binom{n+m}{k}$ by a similar derivation.

11.3 Useful Inequalities

$1 + x \leq e^x$ for all real x .

One often establishes an inequality such as $1 + x \leq e^x$ by showing that the difference of the two sides, namely $e^x - (1 + x)$, is always positive. This can be done by taking derivatives. The first and second derivatives are $e^x - 1$ and e^x . Since e^x is always positive, $e^x - 1$ is monotonic and $e^x - (1 + x)$ is convex. Since $e^x - 1$ is monotonic, it can be zero only once and is zero at $x = 0$. Thus, $e^x - (1 + x)$ takes on its minimum at $x = 0$ where it is zero establishing the inequality.

$$(1 - x)^n \geq 1 - nx \text{ for } 0 \leq x \leq 1$$

Let $g(x) = (1 - x)^n - (1 - nx)$. We establish $g(x) \geq 0$ for x in $[0, 1]$ by taking the derivative.

$$g'(x) = -n(1 - x)^{n-1} + n = n[1 - (1 - x)^{n-1}] \geq 0$$

for $0 \leq x \leq 1$. Thus, g takes on its minimum for x in $[0, 1]$ at $x = 0$ where $g(0) = 0$ proving the inequality.

$$(x + y)^2 \leq 2x^2 + 2y^2$$

The inequality follows from $(x + y)^2 + (x - y)^2 = 2x^2 + 2y^2$.

MOVE THE FOLLOWING TO SECTION "Some Useful Inequalities.

For any non-negative reals $a_1, a_2, \dots, a_n$ and any $\rho \in (0, 1)$, we have $(\sum_{i=1}^n a_i)^\rho \leq \sum_{i=1}^n a_i^\rho$.

Proof We will see that we can reduce to the case when only one of the a_i is non-zero and the rest are zero. To this end, suppose a_1, a_2 are both positive and without loss of generality, assume $a_1 \geq a_2$. Add an infinitesimal positive amount da_1 to a_1 and subtract the same amount from a_2 . This does not alter the left hand side. We claim it does not increase the right hand side: to see this, note that

$$(a_1 + da_1)^\rho + (a_2 - da_1)^\rho - a_1^\rho - a_2^\rho = \rho(a_1^{\rho-1} - a_2^{\rho-1})da_1 + O((da_1)^2),$$

and since $\rho - 1 \leq 0$, we have $a_1^{\rho-1} - a_2^{\rho-1} \leq 0$, proving the claim. Now by repeating this process, we can make $a_2 = 0$ (at that time a_1 will equal the sum of the original a_1 and a_2). Now repeating on all pairs of a_i , we can make all but one of them zero and in the process, we have left the left hand side the same, but have not increased the right hand side. So it suffices to prove the inequality at the end which clearly holds.

$1 + x \leq e^x$ for all real x

$(1 - x)^n \geq 1 - nx$ for $0 \leq x \leq 1$

$(x + y)^2 \leq 2x^2 + 2y^2$

Triangle Inequality $\mathbf{x}$ and $\mathbf{y}$, $|\mathbf{x} + \mathbf{y}| \leq |\mathbf{x}| + |\mathbf{y}|$.

Cauchy-Schwartz Inequality $|\mathbf{x}||\mathbf{y}| \geq \mathbf{x}^T \mathbf{y}$

Young's inequality For positive real numbers p and q where $\frac{1}{p} + \frac{1}{q} = 1$ and positive reals x and y ,

$$xy \leq \frac{1}{p}x^p + \frac{1}{q}y^q.$$

Hölder's inequality For positive real numbers p and q with $\frac{1}{p} + \frac{1}{q} = 1$,

$$\sum_{i=1}^n |x_i y_i| \leq \left(\sum_{i=1}^n |x_i|^p \right)^{1/p} \left(\sum_{i=1}^n |y_i|^q \right)^{1/q}.$$

Jensen's inequality For a convex function f ,

$$f \left(\sum_{i=1}^n \alpha_i x_i \right) \leq \sum_{i=1}^n \alpha_i f(x_i),$$

END OF MOVE TO SECTION "Some useful inequalities.

***** **The Triangle Inequality**

For any two vectors $\mathbf{x}$ and $\mathbf{y}$, $|\mathbf{x} + \mathbf{y}| \leq |\mathbf{x}| + |\mathbf{y}|$. Since $\mathbf{x} \cdot \mathbf{y} \leq |\mathbf{x}||\mathbf{y}|$,

$$|\mathbf{x} + \mathbf{y}|^2 = (\mathbf{x} + \mathbf{y})^T \cdot (\mathbf{x} + \mathbf{y}) = |\mathbf{x}|^2 + |\mathbf{y}|^2 + 2\mathbf{x}^T \cdot \mathbf{y} \leq |\mathbf{x}|^2 + |\mathbf{y}|^2 + 2|\mathbf{x}||\mathbf{y}| = (|\mathbf{x}| + |\mathbf{y}|)^2$$

The inequality follows by taking square roots.

Stirling approximation

$$n! \cong \left(\frac{n}{e}\right)^n \sqrt{2\pi n} \quad \binom{2n}{n} \cong \frac{1}{\sqrt{\pi n}} 2^{2n}$$

$$\sqrt{2\pi n} \frac{n^n}{e^n} < n! < \sqrt{2\pi n} \frac{n^n}{e^n} \left(1 + \frac{1}{12n-1}\right)$$

We prove the inequalities, except for constant factors. Namely, we prove that

$$1.4 \left(\frac{n}{e}\right)^n \sqrt{n} \leq n! \leq e \left(\frac{n}{e}\right)^n \sqrt{n}.$$

Write $\ln(n!) = \ln 1 + \ln 2 + \cdots + \ln n$. This sum is approximately $\int_{x=1}^n \ln x \, dx$. The indefinite integral $\int \ln x \, dx = (x \ln x - x)$ gives an approximation, but without the $\sqrt{n}$ term. To get the $\sqrt{n}$, differentiate twice and note that $\ln x$ is a concave function. This means that for any positive x_0 ,

$$\frac{\ln x_0 + \ln(x_0 + 1)}{2} \leq \int_{x=x_0}^{x_0+1} \ln x \, dx,$$

since for $x \in [x_0, x_0 + 1]$, the curve $\ln x$ is always above the spline joining $(x_0, \ln x_0)$ and $(x_0 + 1, \ln(x_0 + 1))$. Thus,

$$\begin{aligned} \ln(n!) &= \frac{\ln 1}{2} + \frac{\ln 1 + \ln 2}{2} + \frac{\ln 2 + \ln 3}{2} + \cdots + \frac{\ln(n-1) + \ln n}{2} + \frac{\ln n}{2} \\ &\leq \int_{x=1}^n \ln x \, dx + \frac{\ln n}{2} = [x \ln x - x]_1^n + \frac{\ln n}{2} \\ &= n \ln n - n + 1 + \frac{\ln n}{2}. \end{aligned}$$

Thus, $n! \leq n^n e^{-n} \sqrt{n} e$. For the lower bound on $n!$, start with the fact that for any $x_0 \geq 1/2$:

$$\ln x_0 \geq \frac{1}{2}(\ln(x_0 + \rho) + \ln(x_0 - \rho)) \quad \text{implies} \quad \ln x_0 \geq \int_{x=x_0-0.5}^{x_0+0.5} \ln x \, dx.$$

Thus,

$$\ln(n!) = \ln 2 + \ln 3 + \cdots + \ln n \geq \int_{x=1.5}^{n+0.5} \ln x \, dx,$$

from which one can derive a lower bound with a calculation.

Stirling approximation for the binomial coefficient

$$\binom{n}{k} \leq \left(\frac{en}{k}\right)^k$$

Use the Stirling approximation for $k!$:

$$\binom{n}{k} = \frac{n!}{(n-k)!k!} \leq \frac{n^k}{k!} \cong \left(\frac{en}{k}\right)^k.$$

The gamma function

$$\Gamma(a) = \int_0^{\infty} x^{a-1} e^{-x} dx$$

$$\Gamma\left(\frac{1}{2}\right) = \sqrt{\pi}, \quad \Gamma(1) = \Gamma(2) = 1, \text{ and for } n \geq 2, \quad \Gamma(n) = (n-1)\Gamma(n-1).$$

The last statement is proved by induction on n . It is easy to see that $\Gamma(1) = 1$. For $n \geq 2$, we will use integration by parts.

Integration by parts

$$\int f(x) g'(x) dx = f(x) g(x) - \int f'(x) g(x) dx$$

Write $\Gamma(n) = \int_{x=0}^{\infty} f(x)g'(x) dx$, where, $f(x) = x^{n-1}$ and $g'(x) = e^{-x}$. Thus,

$$\Gamma(n) = [f(x)g(x)]_{x=0}^{\infty} + \int_{x=0}^{\infty} (n-1)x^{n-2}e^{-x} dx = (n-1)\Gamma(n-2),$$

as claimed.

Cauchy-Schwartz Inequality

$$\left(\sum_{i=1}^n x_i^2\right) \left(\sum_{i=1}^n y_i^2\right) \geq \left(\sum_{i=1}^n x_i y_i\right)^2$$

In vector form, $|\mathbf{x}||\mathbf{y}| \geq \mathbf{x}^T \mathbf{y}$, the inequality states that the dot product of two vectors is at most the product of their lengths. The Cauchy-Schwartz inequality is a special case of Hölder's inequality below with $p = q = 2$.

Young's inequality

For positive real numbers p and q where $\frac{1}{p} + \frac{1}{q} = 1$ and positive reals x and y ,

$$\frac{1}{p}x^p + \frac{1}{q}y^q \geq xy.$$

The left hand side of Young's inequality, $\frac{1}{p}x^p + \frac{1}{q}y^q$, is a convex combination of x^p and y^q since $\frac{1}{p}$ and $\frac{1}{q}$ sum to 1. Since the $\ln$ of a convex combination of two elements is greater than or equal to the convex combination of the $\ln$ of the elements

$$\ln\left(\frac{1}{p}x^p + \frac{1}{q}y^q\right) \geq \frac{1}{p}\ln(x^p) + \frac{1}{q}\ln(y^q) = \ln(xy).$$

Since for $x \geq 0$, $\ln x$ is a monotone increasing function, $\frac{1}{p}x^p + \frac{1}{q}y^q \geq xy$.

Hölder's inequality

For any positive real numbers p and q with $\frac{1}{p} + \frac{1}{q} = 1$,

$$\sum_{i=1}^n |x_i y_i| \leq \left(\sum_{i=1}^n |x_i|^p \right)^{1/p} \left(\sum_{i=1}^n |y_i|^q \right)^{1/q}.$$

Let $x'_i = x_i / (\sum_{i=1}^n |x_i|^p)^{1/p}$ and $y'_i = y_i / (\sum_{i=1}^n |y_i|^q)^{1/q}$. Replacing x_i by x'_i and y_i by y'_i does not change the inequality. Now $\sum_{i=1}^n |x'_i|^p = \sum_{i=1}^n |y'_i|^q = 1$, so it suffices to prove $\sum_{i=1}^n |x'_i y'_i| \leq 1$. Apply Young's inequality to get $|x'_i y'_i| \leq \frac{|x'_i|^p}{p} + \frac{|y'_i|^q}{q}$. Summing over i , the right hand side sums to $\frac{1}{p} + \frac{1}{q} = 1$ finishing the proof.

For $a_1, a_2, \dots, a_n$ real and k a positive integer,

$$(a_1 + a_2 + \dots + a_n)^k \leq n^{k-1}(|a_1|^k + |a_2|^k + \dots + |a_n|^k).$$

Using Hölder's inequality with $p = k$ and $q = k/(k-1)$,

$$\begin{aligned} |a_1 + a_2 + \dots + a_n| &\leq |a_1 \cdot 1| + |a_2 \cdot 1| + \dots + |a_n \cdot 1| \\ &\leq \left(\sum_{i=1}^n |a_i|^k \right)^{1/k} (1 + 1 + \dots + 1)^{(k-1)/k}, \end{aligned}$$

from which the current inequality follows.

Arithmetic and geometric means

The arithmetic mean of a set of nonnegative reals is at least their geometric mean. For $a_1, a_2, \dots, a_n > 0$

$$\frac{1}{n} \sum_{i=1}^n a_i \geq \sqrt[n]{a_1 a_2 \dots a_n}$$

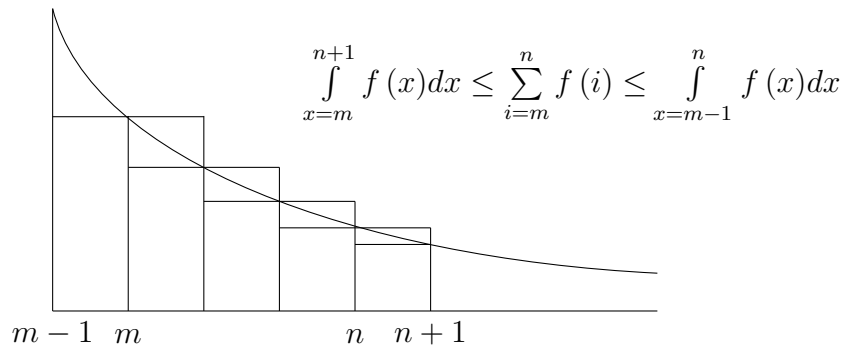


Figure 11.1: Approximating sums by integrals

Assume that $a_1 \geq a_2 \geq \dots \geq a_n$. We reduce the proof to the case when all the a_i are equal; in which case the inequality holds with equality. Suppose $a_1 > a_2$. Let ε be a positive infinitesimal. Add ε to a_2 and subtract ε from a_1 (to get closer to the case when they are equal). The left hand side $\frac{1}{n} \sum_{i=1}^n a_i$ does not change.

$$\begin{aligned} (a_1 - \varepsilon)(a_2 + \varepsilon)a_3a_4 \cdots a_n &= a_1a_2 \cdots a_n + \varepsilon(a_1 - a_2)a_3a_4 \cdots a_n + O(\varepsilon^2) \\ &> a_1a_2 \cdots a_n \end{aligned}$$

for small enough $\varepsilon > 0$. Thus, the change has increased $\sqrt[n]{a_1a_2 \cdots a_n}$. So if the inequality holds after the change, it must hold before. By continuing this process, one can make all the a_i equal.

Approximating sums by integrals

For monotonic decreasing $f(x)$,

$$\int_{x=m}^{n+1} f(x) dx \leq \sum_{i=m}^n f(i) \leq \int_{x=m-1}^n f(x) dx.$$

See Fig. 11.1. Thus,

$$\int_{x=2}^{n+1} \frac{1}{x^2} dx \leq \sum_{i=2}^n \frac{1}{i^2} = \frac{1}{4} + \frac{1}{9} + \cdots + \frac{1}{n^2} \leq \int_{x=1}^n \frac{1}{x^2} dx$$

and hence $\frac{3}{2} - \frac{1}{n+1} \leq \sum_{i=1}^n \frac{1}{i^2} \leq 2 - \frac{1}{n}$.

Jensen's Inequality

For a convex function f ,

$$f\left(\frac{1}{2}(x_1 + x_2)\right) \leq \frac{1}{2}(f(x_1) + f(x_2)).$$

More generally for any convex function f ,

$$f\left(\sum_{i=1}^n \alpha_i x_i\right) \leq \sum_{i=1}^n \alpha_i f(x_i),$$

where $0 \leq \alpha_i \leq 1$ and $\sum_{i=1}^n \alpha_i = 1$. From this, it follows that for any convex function f and random variable x ,

$$E(f(x)) \geq f(E(x)).$$

We prove this for a discrete random variable x taking on values $a_1, a_2, \dots$ with $\text{Prob}(x = a_i) = \alpha_i$:

$$E(f(x)) = \sum_i \alpha_i f(a_i) \geq f\left(\sum_i \alpha_i a_i\right) = f(E(x)).$$

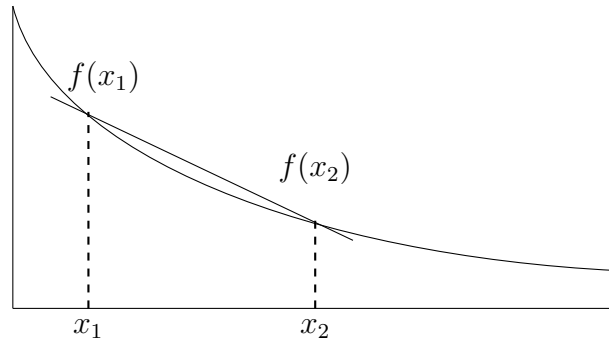


Figure 11.2: For a convex function f , $f\left(\frac{x_1+x_2}{2}\right) \leq \frac{1}{2}(f(x_1) + f(x_2))$.

Example: Let $f(x) = x^k$ for k an even positive integer. Then, $f''(x) = k(k-1)x^{k-2}$ which since $k-2$ is even is nonnegative for all x implying that f is convex. Thus,

$$E(x) \leq \sqrt[k]{E(x^k)},$$

since $t^{\frac{1}{k}}$ is a monotone function of t , $t > 0$. It is easy to see that this inequality does not necessarily hold when k is odd; indeed for odd k , x^k is not a convex function. ■

Tails of Gaussian

For bounding the tails of Gaussian densities, the following inequality is useful. The proof uses a technique useful in many contexts. For $t > 0$,

$$\int_{x=t}^{\infty} e^{-x^2} dx \leq \frac{e^{-t^2}}{2t}.$$

In proof, first write: $\int_{x=t}^{\infty} e^{-x^2} dx \leq \int_{x=t}^{\infty} \frac{x}{t} e^{-x^2} dx$, using the fact that $x \geq t$ in the range of integration. The latter expression is integrable in closed form since $d(e^{-x^2}) = (-2x)e^{-x^2}$ yielding the claimed bound.

A similar technique yields an upper bound on

$$\int_{x=\beta}^1 (1-x^2)^\alpha dx,$$

for $\beta \in [0, 1]$ and $\alpha > 0$. Just use $(1-x^2)^\alpha \leq \frac{x}{\beta}(1-x^2)^\alpha$ over the range and integrate in closed form the last expression.

$$\begin{aligned} \int_{x=\beta}^1 (1-x^2)^\alpha dx &\leq \int_{x=\beta}^1 \frac{x}{\beta} (1-x^2)^\alpha dx = \frac{-1}{2\beta(\alpha+1)} (1-x^2)^{\alpha+1} \Big|_{x=\beta}^1 \\ &= \frac{(1-\beta^2)^{\alpha+1}}{2\beta(\alpha+1)} \end{aligned}$$

11.4 Probability

Consider an experiment such as flipping a coin whose outcome is determined by chance. To talk about the outcome of a particular experiment, we introduce the notion of a *random variable* whose value is the outcome of the experiment. The set of possible outcomes is called the *sample space*. If the sample space is finite, we can assign a probability of occurrence to each outcome. In some situations where the sample space is infinite, we can assign a probability of occurrence. The probability $p(i) = \frac{6}{\pi^2} \frac{1}{i^2}$ for i an integer greater than or equal to one is such an example. The function assigning the probabilities is called a *probability distribution function*.

In many situations, a probability distribution function does not exist. For example, for the uniform probability on the interval $[0,1]$, the probability of any specific value is

zero. What we can do is define a *probability density function* $p(x)$ such that

$$\text{Prob}(a < x < b) = \int_a^b p(x)dx$$

If x is a continuous random variable for which a density function exists, then the *cumulative distribution function* $f(a)$ is defined by

$$f(a) = \int_{-\infty}^a p(x)dx$$

which gives the probability that $x \leq a$.

11.4.1 Sample Space, Events, Independence

There may be more than one relevant random variable in a situation. For example, if one tosses n coins, there are n random variables, $x_1, x_2, \dots, x_n$, taking on values 0 and 1, a 1 for heads and a 0 for tails. The set of possible outcomes, the sample space, is $\{0, 1\}^n$. An *event* is a subset of the sample space. The event of an odd number of heads, consists of all elements of $\{0, 1\}^n$ with an odd number of 1's.

For two events A and B , the joint probability is denoted $\text{Prob}(A, B)$. The *conditional probability* of A given that B has occurred is denoted by $\text{Prob}(A|B)$ and is given by

$$\text{Prob}(A|B) = \frac{\text{Prob}(A, B)}{\text{Prob}(B)}.$$

Events A and B are *independent* if the occurrence of one event has no influence on the probability of the other. That is, $\text{Prob}(A|B) = \text{Prob}(A)$ or equivalently, $\text{Prob}(A, B) = \text{Prob}(A)\text{Prob}(B)$. Two random variables x and y are *independent* if for every possible set A of values of x and every possible set B of values of y , the events A and B are independent.

If (x, y) is a random vector and one normalizes it to a unit vector $\left(\frac{x}{x^2+y^2}, \frac{y}{x^2+y^2}\right)$ the coordinates are no longer independent since knowing the value of one coordinate uniquely determines the value of the other.

A collection of n random variables $x_1, x_2, \dots, x_n$ is *mutually independent* if for all possible sets $A_1, A_2, \dots, A_n$ of values of $x_1, x_2, \dots, x_n$,

$$\text{Prob}(A_1, A_2, \dots, A_n) = \text{Prob}(A_1)\text{Prob}(A_2) \cdots \text{Prob}(A_n).$$

If the random variables are discrete, it would suffice to say that for any real numbers $a_1, a_2, \dots, a_n$

$$\text{Prob}(x_1 = a_1, x_2 = a_2, \dots, x_n = a_n) = \text{Prob}(x_1 = a_1)\text{Prob}(x_2 = a_2) \cdots \text{Prob}(x_n = a_n).$$

Mutual independence is much stronger than requiring that the variables are pairwise independent. Consider the example of 2-universal hash functions discussed in Chapter ??.

11.4.2 Linearity of Expectation

An important concept is that of the expectation of a random variable. The *expected value*, $E(x)$, of a random variable x is $E(x) = \sum_x xp(x)$ in the discrete case and $E(x) = \int_{-\infty}^{\infty} xp(x)dx$ in the continuous case. The expectation of a sum of random variables is equal to the sum of their expectations. The linearity of expectation follows directly from the definition and does not require independence.

11.4.3 Indicator Variables

A useful tool is that of an indicator variable that takes on value 0 or 1 to indicate whether some quantity is present or not. The indicator variable is useful in determining the expected size of a subset. Given a random subset of the integers $\{1, 2, \dots, n\}$, the expected size of the subset is the expected value of $x_1 + x_2 + \dots + x_n$ where x_i is the indicator variable that takes on value 1 if i is in the subset.

Example: Consider a random permutation of n integers. Define the indicator function $x_i = 1$ if the i^{th} integer in the permutation is i . The expected number of fixed points is given by

$$E\left(\sum_{i=1}^n x_i\right) = \sum_{i=1}^n E(x_i) = n \frac{1}{n} = 1.$$

Note that the x_i are not independent. But, linearity of expectation still applies. ■

Example: Consider the expected number of vertices of degree d in a random graph $G(n, p)$. The number of vertices of degree d is the sum of n indicator random variables, one for each vertex, with value one if the vertex has degree d . The expectation is the sum of the expectations of the n indicator random variables and this is just n times the expectation of one of them. ■

11.4.4 Variance

In addition to the expected value of a random variable, another important parameter is the variance. The *variance* of a random variable x , often denoted $\sigma^2(x)$ is $E(x - E(x))^2$ and measures how close to the expected value the random variable is likely to be. The *standard deviation* σ is the square root of the variance. The units of σ are the same as

those of x .

By linearity of expectation

$$\sigma^2 = E(x - E(x))^2 = E(x^2) - 2E(x)E(x) + E^2(x) = E(x^2) - E^2(x).$$

11.4.5 Variance of the Sum of Independent Random Variables

In general, the variance of the sum is not equal to the sum of the variances. However, if x and y are independent, then $E(xy) = E(x)E(y)$ and

$$\text{var}(x + y) = \text{var}(x) + \text{var}(y).$$

To see this

$$\begin{aligned} \text{var}(x + y) &= E((x + y)^2) - E^2(x + y) \\ &= E(x^2) + 2E(xy) + E(y^2) - E^2(x) - 2E(x)E(y) - E^2(y) \end{aligned}$$

From independence $2E(xy) - 2E(x)E(y) = 0$ and

$$\begin{aligned} \text{var}(x + y) &= E(x^2) - E^2(x) + E(y^2) - E^2(y) \\ &= \text{var}(x) + \text{var}(y) \end{aligned}$$

More generally, if $x_1, x_2, \dots, x_n$ are pairwise independent random variables, then

$$\text{var}(x_1 + x_2 + \dots + x_n) = \text{var}(x_1) + \text{var}(x_2) + \dots + \text{var}(x_n).$$

For the variance of the sum to be the sum of the variances only requires pairwise independence not full independence.

11.4.6 The Central Limit Theorem

Let $s = x_1 + x_2 + \dots + x_n$ be a sum of n independent random variables where each x_i has probability distribution

$$x_i = \begin{cases} 0 & \frac{1}{2} \\ 1 & \frac{1}{2} \end{cases}.$$

The expected value of each x_i is $1/2$ with variance

$$\sigma_i^2 = \left(\frac{1}{2} - 0\right)^2 \frac{1}{2} + \left(\frac{1}{2} - 1\right)^2 \frac{1}{2} = \frac{1}{4}.$$

The expected value of s is $n/2$ and since the variables are independent, the variance of the sum is the sum of the variances and hence is $n/4$. How concentrated s is around its mean depends on the standard deviation of s which is $\frac{\sqrt{n}}{2}$. For n equal 100 the expected value of s is 50 with a standard deviation of 5 which is 10% of the mean. For $n = 10,000$ the expected value of s is 5,000 with a standard deviation of 50 which is 1% of the mean. Note that as n increases, the standard deviation increases, but the ratio of the standard deviation to the mean goes to zero. The central limit theorem quantifies this.

Theorem 11.1 *Suppose $x_1, x_2, \dots, x_n$ is a sequence of independent identically distributed random variables, each with mean μ and variance σ^2 . The distribution of the random variable*

$$\frac{1}{\sqrt{n}}(x_1 + x_2 + \dots + x_n - n\mu)$$

converges to the distribution of the Gaussian with mean 0 and variance σ^2 .

The variance of $x_1 + x_2 + \dots + x_n$ behaves as $n\sigma^2$ since it is the sum of n independent random variables, each with variance σ^2 . Dividing the sum $x_1 + x_2 + \dots + x_n$ by $\sqrt{n}$ makes the variance bounded. Recall that the limiting value has to be bounded for the limit to exist.

11.4.7 Median

One often calculates the average value of a random variable to get a feeling for the magnitude of the variable. This is reasonable when the probability distribution of the variable is Gaussian, or has a small variance. However, if there are outliers or if the probability is not concentrated about its expected value, then the average may be distorted by outliers. An alternative to calculating the expected value is to calculate the median, the value for which half of the probability is above and half is below.

11.4.8 Unbiased Estimators

Consider n samples $x_1, x_2, \dots, x_n$ from a Gaussian distribution of mean μ and variance σ^2 . For this distribution $m = \frac{x_1 + x_2 + \dots + x_n}{n}$ is an unbiased estimator of μ , which means that $E(m) = \mu$ and $\frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2$ is an unbiased estimator of σ^2 . However, if μ is not known and is approximated by m , then $\frac{1}{n-1} \sum_{i=1}^n (x_i - m)^2$ is an unbiased estimator of σ^2 .

11.4.9 Probability Distributions

The Gaussian or normal distribution

The normal distribution is

$$\frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2} \frac{(x-m)^2}{\sigma^2}}$$

where m is the mean and σ^2 is the variance. The coefficient $\frac{1}{\sqrt{2\pi}\sigma}$ makes the integral of the distribution be one. If we measure distance in units of the standard deviation σ from the mean, then

$$\phi(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}x^2}$$

Standard tables give values of the integral

$$\int_0^t \phi(x) dx$$

and from these values one can compute probability integrals for a normal distribution with mean m and variance σ^2 .

Bernoulli trials and the binomial distribution

A Bernoulli trial has two possible outcomes, called success or failure, with probabilities p and $1 - p$, respectively. If there are n independent Bernoulli trials, the probability of exactly k successes is given by the *binomial distribution*

$$B(n, p) = \binom{n}{k} p^k (1 - p)^{n-k}$$

The mean and variance of the binomial distribution $B(n, p)$ are np and $np(1 - p)$, respectively. The mean of the binomial distribution is np , by linearity of expectations. The variance is $np(1 - p)$ since the variance of a sum of independent random variables is the sum of their variances.

Let x_1 be the number of successes in n_1 trials and let x_2 be the number of successes in n_2 trials. The probability distribution of the sum of the successes, $x_1 + x_2$, is the same as the distribution of $x_1 + x_2$ successes in $n_1 + n_2$ trials. Thus, $B(n_1, p) + B(n_2, p) = B(n_1 + n_2, p)$.

Poisson distribution

The Poisson distribution describes the probability of k events happening in a unit of time when the average rate per unit of time is λ . Divide the unit of time into n segments. When n is large enough, each segment is sufficiently small so that the probability of two events happening in the same segment is negligible. The Poisson distribution gives the probability of k events happening in a unit of time and can be derived from the binomial distribution by taking the limit as $n \rightarrow \infty$.

Let $p = \frac{\lambda}{n}$. Then

$$\begin{aligned} \text{Prob}(k \text{ successes in a unit of time}) &= \lim_{n \rightarrow \infty} \binom{n}{k} \left(\frac{\lambda}{n}\right)^k \left(1 - \frac{\lambda}{n}\right)^{n-k} \\ &= \lim_{n \rightarrow \infty} \frac{n(n-1) \cdots (n-k+1)}{k!} \left(\frac{\lambda}{n}\right)^k \left(1 - \frac{\lambda}{n}\right)^n \left(1 - \frac{\lambda}{n}\right)^{-k} \\ &= \lim_{n \rightarrow \infty} \frac{\lambda^k}{k!} e^{-\lambda} \end{aligned}$$

In the limit as n goes to infinity the binomial distribution $p(k) = \binom{n}{k} p^k (1-p)^{n-k}$ becomes the Poisson distribution $p(k) = e^{-\lambda} \frac{\lambda^k}{k!}$. The mean and the variance of the Poisson distribution have value λ . If x and y are both Poisson random variables from distributions with means λ_1 and λ_2 respectively, then $x + y$ is Poisson with mean $m_1 + m_2$. For large n and small p the binomial distribution can be approximated with the Poisson distribution.

The binomial distribution with mean np and variance $np(1-p)$ can be approximated by the normal distribution with mean np and variance $np(1-p)$. The central limit theorem tells us that there is such an approximation in the limit. The approximation is good if both np and $n(1-p)$ are greater than 10 provided k is not extreme. Thus,

$$\binom{n}{k} \left(\frac{1}{2}\right)^k \left(\frac{1}{2}\right)^{n-k} \cong \frac{1}{\sqrt{\pi n/2}} e^{-\frac{(n/2-k)^2}{\frac{1}{2}n}}.$$

This approximation is excellent provided k is $\Theta(n)$. The Poisson approximation

$$\binom{n}{k} p^k (1-p)^{n-k} \cong e^{-np} \frac{(np)^k}{k!}$$

is off for central values and tail values even for $p = 1/2$. The approximation

$$\binom{n}{k} p^k (1-p)^{n-k} \cong \frac{1}{\sqrt{\pi p n}} e^{-\frac{(pn-k)^2}{pn}}$$

is good for $p = 1/2$ but error for other values of p .

Generation of random numbers according to a given probability distribution

Suppose one wanted to generate a random variable with probability density $p(x)$ where $p(x)$ is continuous. Let $P(x)$ be the cumulative distribution function for x and let u be a random variable with uniform probability density over the interval $[0,1]$. Then the random variable $x = P^{-1}(u)$ has probability density $p(x)$.

Example: For a Cauchy density function the cumulative distribution function is

$$P(x) = \int_{t=-\infty}^x \frac{1}{\pi} \frac{1}{1+t^2} dt = \frac{1}{2} + \frac{1}{\pi} \tan^{-1}(x).$$

Setting $u = P(x)$ and solving for x yields $x = \tan\left(\pi\left(u - \frac{1}{2}\right)\right)$. Thus, to generate a random number $x \geq 0$ using the Cauchy distribution, generate u , $0 \leq u \leq 1$, uniformly and calculate $x = \tan\left(\pi\left(u - \frac{1}{2}\right)\right)$. The value of x varies from $-\infty$ to ∞ with $x = 0$ for $u = 1/2$. ■

11.4.10 Maximum Likelihood Estimation MLE

Suppose the probability distribution of a random variable x depends on a parameter r . With slight abuse of notation since r is a parameter rather than a random variable, we denote the probability distribution of x as $p(x|r)$. This is the likelihood of observing x if r was in fact the parameter value. The job of the maximum likelihood estimator, MLE, is to find the best r after observing values of the random variable x . The likelihood of r given the observed values of x , denoted $L(r|x)$ is the probability of observing the observed value of x as a function of the parameter r . The maximum likelihood is the value of r that maximizes the function $L(r|x)$.

Example: Consider flipping a coin 100 times. Suppose 62 heads and 38 tails occur. What is the most likely value of the probability of the coin to come down heads when the coin is flipped? In this case, it is $r = 0.62$. The probability that we get 62 heads if the unknown probability of heads in one trial is r is

$$\text{Prob}(62 \text{ heads}|r) = \binom{100}{62} r^{62} (1-r)^{38}.$$

This quantity is maximized when $r = 0.62$. To see this take the logarithm, which as a function of r is $\ln\left(\binom{100}{62}\right) + 62 \ln r + 38 \ln(1-r)$. The derivative with respect to r is zero at $r = 0.62$ and the second derivative is negative indicating a maximum. Thus, $r = 0.62$ is the maximum likelihood estimator of the probability of heads in a trial. ■

Bayes rule

Given a joint probability distribution $\text{Prob}(A, B)$, Bayes rule relates the conditional probability of A given B to the conditional probability of B given A .

$$\text{Prob}(A|B) = \frac{\text{Prob}(B|A) \text{Prob}(A)}{\text{Prob}(B)}$$

Suppose one knows the probability of A and wants to know how this probability changes if we know that B has occurred. $\text{Prob}(A)$ is called the prior probability. The conditional

probability $\text{Prob}(A|B)$ is called the posterior probability because it is the probability of A after we know that B has occurred.

The example below illustrates that if a situation is rare, a highly accurate test will often give the wrong answer.

Example: Let A be the event that a product is defective and let B be the event that a test says a product is defective. Let $\text{Prob}(B|A)$ be the probability that the test says a product is defective assuming the product is defective and let $\text{Prob}(B|\bar{A})$ be the probability that the test says a product is defective if it is not actually defective.

What is the probability $\text{Prob}(A|B)$ that the product is defective if the test say it is defective? Suppose $\text{Prob}(A) = 0.001$, $\text{Prob}(B|A) = 0.99$, and $\text{Prob}(B|\bar{A}) = 0.02$. Then

$$\begin{aligned}\text{Prob}(B) &= \text{Prob}(B|A)\text{Prob}(A) + \text{Prob}(B|\bar{A})\text{Prob}(\bar{A}) \\ &= 0.99 \times 0.001 + 0.02 \times 0.999 \\ &= 0.02087\end{aligned}$$

and

$$\text{Prob}(A|B) = \frac{\text{Prob}(B|A)\text{Prob}(A)}{\text{Prob}(B)} \approx \frac{0.99 \times 0.001}{0.0210} = 0.0471$$

Even though the test fails to detect a defective product only 1% of the time when it is defective and claims that it is defective when it is not only 2% of the time, the test is correct only 4.7% of the time when it says a product is defective. This comes about because of the low frequencies of defective products. ■

The words prior, a posteriori, and likelihood come from Bayes theorem.

$$\text{a posteriori} = \frac{\text{likelihood} \times \text{prior}}{\text{normalizing constant}}$$

$$\text{Prob}(A|B) = \frac{\text{Prob}(B|A)\text{Prob}(A)}{\text{Prob}(B)}$$

The a posteriori probability is the conditional probability of A given B . The likelihood is the conditional probability $\text{Prob}(B|A)$.

11.4.11 Tail Bounds

Markov's inequality bounds the probability that a nonnegative random variable exceeds a value a .

$$p(x \geq a) \leq \frac{E(x)}{a}.$$

or

$$p(x \geq aE(x)) \leq \frac{1}{a}$$

If one also knows the variance, σ^2 , then using Chebycheff's inequality one can bound the probability that a random variable differs from its expected value by more than a standard deviations.

$$p(|x - m| \geq a\sigma) \leq \frac{1}{a^2}$$

If a random variable s is the sum of n independent random variables $x_1, x_2, \dots, x_n$ of finite variance, then better bounds are possible. For any $\delta > 0$,

$$\text{Prob}(s > (1 + \delta)m) < \left[\frac{e^\delta}{(1 + \delta)^{(1+\delta)}} \right]^m$$

and for $0 < \gamma \leq 1$,

$$\text{Prob}(s < (1 - \gamma)m) < \left[\frac{e^{-\gamma}}{(1 + \gamma)^{(1+\gamma)}} \right]^m < e^{-\frac{\gamma^2 m}{2}}$$

11.5 Chernoff Bounds

Chebyshev's inequality bounds the probability that a random variable will deviate from its mean by more than a given amount. Chebyshev's inequality holds for any probability distribution. For some distributions we can get much tighter bounds. For example, the probability that a Gaussian random variable deviates from its mean falls off exponentially with the distance from the mean. Here we shall be concerned with the situation where we have a random variable that is the sum of n independent random variables. This is another situation in which we can derive a tighter bound than that given by the Chebyshev inequality. We consider the case where the n independent variables are binomial but similar results can be shown for independent random variables from any distribution that has a finite variance.

Let $x_1, x_2, \dots, x_n$ be independent random variables where

$$x_i = \begin{cases} 0 & \text{Prob } 1 - p \\ 1 & \text{Prob } p \end{cases}.$$

Consider the sum $s = \sum_{i=1}^n x_i$. Here the expected value of each x_i is p and by linearity of expectation, the expected value of the sum is $m = np$. Theorem ?? bounds the probability that the sum s exceeds $(1 + \delta)m$.

Theorem 11.2 For any $\delta > 0$, $\text{Prob}(s > (1 + \delta)m) < \left(\frac{e^\delta}{(1 + \delta)^{(1+\delta)}} \right)^m$

Proof: For any $\lambda > 0$, the function $e^{\lambda x}$ is monotone. Thus,

$$\text{Prob}(s > (1 + \delta)m) = \text{Prob}(e^{\lambda s} > e^{\lambda(1+\delta)m}).$$

$e^{\lambda x}$ is nonnegative for all x , so we can apply Markov's inequality to get

$$\text{Prob}(e^{\lambda s} > e^{\lambda(1+\delta)m}) \leq e^{-\lambda(1+\delta)m} E(e^{\lambda s}).$$

Since the x_i are independent,

$$\begin{aligned} E(e^{\lambda s}) &= E\left(e^{\lambda \sum_{i=1}^n x_i}\right) = E\left(\prod_{i=1}^n e^{\lambda x_i}\right) = \prod_{i=1}^n E(e^{\lambda x_i}) \\ &= \prod_{i=1}^n (e^{\lambda p} + 1 - p) = \prod_{i=1}^n (p(e^{\lambda} - 1) + 1). \end{aligned}$$

Using the inequality $1 + x < e^x$ with $x = p(e^{\lambda} - 1)$ yields

$$E(e^{\lambda s}) < \prod_{i=1}^n e^{p(e^{\lambda} - 1)}.$$

Thus, for all $\lambda > 0$

$$\begin{aligned} \text{Prob}(s > (1 + \delta)m) &\leq \text{Prob}(e^{\lambda s} > e^{\lambda(1+\delta)m}) \\ &\leq e^{-\lambda(1+\delta)m} E(e^{\lambda s}) \\ &\leq e^{-\lambda(1+\delta)m} \prod_{i=1}^n e^{p(e^{\lambda} - 1)}. \end{aligned}$$

Setting $\lambda = \ln(1 + \delta)$

$$\begin{aligned} \text{Prob}(s > (1 + \delta)m) &\leq (e^{-\ln(1+\delta)})^{(1+\delta)m} \prod_{i=1}^n e^{p(e^{\ln(1+\delta)} - 1)} \\ &\leq \left(\frac{1}{1 + \delta}\right)^{(1+\delta)m} \prod_{i=1}^n e^{p\delta} \\ &\leq \left(\frac{1}{(1 + \delta)}\right)^{(1+\delta)m} e^{np\delta} \\ &\leq \left(\frac{e^{\delta}}{(1 + \delta)^{(1+\delta)}}\right)^m. \end{aligned}$$

■

To simplify the bound of Theorem ??, observe that

$$(1 + \delta) \ln(1 + \delta) = \delta + \frac{\delta^2}{2} - \frac{\delta^3}{6} + \frac{\delta^4}{12} - \dots.$$

Therefore

$$(1 + \delta)^{(1+\delta)} = e^{\delta + \frac{\delta^2}{2} - \frac{\delta^3}{6} + \frac{\delta^4}{12} - \dots}$$

and hence

$$\frac{e^\delta}{(1+\delta)^{(1+\delta)}} = e^{-\frac{\delta^2}{2} + \frac{\delta^3}{6} - \dots}.$$

Thus, the bound simplifies to

$$\text{Prob}(s < (1 + \delta)m) \leq e^{-\frac{\delta^2}{2}m + \frac{\delta^3}{6}m - \dots}.$$

For small δ the probability drops exponentially with δ^2 .

When δ is large another simplification is possible. First

$$\text{Prob}(s > (1 + \delta)m) \leq \left(\frac{e^\delta}{(1 + \delta)^{(1+\delta)}} \right)^m \leq \left(\frac{e}{1 + \delta} \right)^{(1+\delta)m}$$

If $\delta > 2e - 1$, substituting $2e - 1$ for δ in the denominator yields

$$\text{Prob}(s > (1 + \delta)m) \leq 2^{-(1+\delta)m}.$$

Theorem ?? gives a bound on the probability of the sum being greater than the mean. We now bound the probability that the sum will be less than its mean.

Theorem 11.3 *Let $0 < \gamma \leq 1$, then $\text{Prob}(s < (1 - \gamma)m) < \left(\frac{e^{-\gamma}}{(1+\gamma)^{(1+\gamma)}} \right)^m < e^{-\frac{\gamma^2 m}{2}}$.*

Proof: For any $\lambda > 0$

$$\text{Prob}(s < (1 - \gamma)m) = \text{Prob}(-s > -(1 - \gamma)m) = \text{Prob}(e^{-\lambda s} > e^{-\lambda(1-\gamma)m}).$$

Applying Markov's inequality

$$\text{Prob}(s < (1 - \gamma)m) < \frac{E(e^{-\lambda s})}{e^{-\lambda(1-\gamma)m}} < \frac{\prod_{i=1}^n E(e^{-\lambda X_i})}{e^{-\lambda(1-\gamma)m}}.$$

Now

$$E(e^{-\lambda x_i}) = pe^{-\lambda} + 1 - p = 1 + p(e^{-\lambda} - 1) + 1.$$

Thus,

$$\text{Prob}(s < (1 - \gamma)m) < \frac{\prod_{i=1}^n [1 + p(e^{-\lambda} - 1)]}{e^{-\lambda(1-\gamma)m}}.$$

Since $1 + x < e^x$

$$\text{Prob}(s < (1 - \gamma)m) < \frac{e^{np(e^{-\lambda}-1)}}{e^{-\lambda(1-\gamma)m}}.$$

Setting $\lambda = \ln \frac{1}{1-\gamma}$

$$\begin{aligned} \text{Prob}(s < (1 - \gamma)m) &< \frac{e^{np(1-\gamma-1)}}{(1 - \gamma)^{(1-\gamma)m}} \\ &< \left(\frac{e^{-\gamma}}{(1 - \gamma)^{(1-\gamma)}} \right)^m. \end{aligned}$$

But for $0 < \gamma \leq 1$, $(1 - \gamma)^{(1-\gamma)} > e^{-\gamma + \frac{\gamma^2}{2}}$. To see this note that

$$\begin{aligned} (1 - \gamma) \ln(1 - \gamma) &= (1 - \gamma) \left(-\gamma - \frac{\gamma^2}{2} - \frac{\gamma^3}{3} - \dots \right) \\ &= -\gamma - \frac{\gamma^2}{2} - \frac{\gamma^3}{3} - \dots + \gamma^2 + \frac{\gamma^3}{2} + \frac{\gamma^4}{3} + \dots \\ &= -\gamma + \left(\gamma^2 - \frac{\gamma^2}{2} \right) + \left(\frac{\gamma^3}{2} - \frac{\gamma^3}{3} \right) + \dots \\ &= -\gamma + \frac{\gamma^2}{2} + \frac{\gamma^3}{6} + \dots \\ &\geq -\gamma + \frac{\gamma^2}{2}. \end{aligned}$$

It then follows that

$$\text{Prob}(s < (1 - \gamma)m) < \left(\frac{e^{-\gamma}}{(1 - \gamma)^{(1-\gamma)}} \right)^m < e^{-\frac{m\gamma^2}{2}}.$$

■

Chernoff's inequality bounds the probability of the deviation of a sum of binomial random variables from its mean. Hoeffding (1963) extends this result to sums of bounded random variables.

Theorem 11.4 (*Devroye et al p122*)(*Hoeffding(1963)*). *Let $x_1, x_2, \dots, x_n$ be independent bounded random variables such that each x_i falls in an interval $[a_i, b_i]$ with probability one.*

Let $S_n = \sum_{i=1}^n x_i$ and $l^2 = \sum_{i=1}^n (b_i - a_i)^2$. For any $\varepsilon > 0$

$$\begin{aligned} \text{Prob}[(S_n - E(S_n)) \geq \varepsilon] &\leq e^{-\frac{2\varepsilon^2}{l^2}} \\ \text{Prob}[(S_n - E(S_n)) \leq -\varepsilon] &\leq e^{-\frac{2\varepsilon^2}{l^2}} \end{aligned}$$

Proof: See Devroye et al pp122-123 ■

Example: In n trials each with probability p of success, what is the probability that the total number of successes x exceeds the expected number, pn , by more than ε . Hoeffding's bound states

$$\text{Prob}(x - pn \geq \varepsilon) \leq e^{-\frac{2\varepsilon^2}{n}}.$$

Thus, if $p = \frac{1}{2}$, the probability of exceeding $0.5n$ by more than say $\sqrt{n}$ is less than e^{-2} .

If $p = \frac{1}{n}$, the probability of exceeding one by $\log n$ is less than $e^{-\frac{2 \log^2 n}{n}}$, which goes asymptotically to one and the bound provides no information. ■

11.6 Eigenvalues and Eigenvectors

11.6.1 Eigenvalues and Eigenvectors

Let A be an $n \times n$ real matrix. The scalar λ is called an eigenvalue of A if there exists a nonzero vector $\mathbf{x}$ satisfying the equation $A\mathbf{x} = \lambda\mathbf{x}$. The vector $\mathbf{x}$ is called the eigenvector of A associated with λ . The set of all eigenvectors associated with a given eigenvalue form a subspace as seen from the fact that if $A\mathbf{x} = \lambda\mathbf{x}$ and $A\mathbf{y} = \lambda\mathbf{y}$, then for any scalars c and d , $A(c\mathbf{x} + d\mathbf{y}) = \lambda(c\mathbf{x} + d\mathbf{y})$. The equation $A\mathbf{x} = \lambda\mathbf{x}$ has a nontrivial solution only if $\det(A - \lambda I) = 0$. The equation $\det(A - \lambda I) = 0$ is called the *characteristic* equation and has n not necessarily distinct roots.

Matrices A and B are similar if there is an invertible matrix P such that $A = P^{-1}BP$.

Theorem 11.5 *If A and B are similar, then they have the same eigenvalues.*

Proof: Let A and B be similar matrices. Then there exists an invertible matrix P such that $A = P^{-1}BP$. For an eigenvector $\mathbf{x}$ of A with eigenvalue λ , $A\mathbf{x} = \lambda\mathbf{x}$, which implies $P^{-1}BP\mathbf{x} = \lambda\mathbf{x}$ or $B(P\mathbf{x}) = \lambda(P\mathbf{x})$. So, $P\mathbf{x}$ is an eigenvector of B with the same eigenvalue λ . Since the reverse also holds, the theorem follows. ■

Even though two similar matrices, A and B , have the same eigenvalues, their eigenvectors are in general different.

The matrix A is *diagonalizable* if A is similar to a diagonal matrix.

Theorem 11.6 *A is diagonalizable if and only if A has n linearly independent eigenvectors.*

Proof:

(only if) Assume A is diagonalizable. Then there exists an invertible matrix P and a diagonal matrix D such that $D = P^{-1}AP$. Thus $PD = AP$. Let the diagonal elements of D be $\lambda_1, \lambda_2, \dots, \lambda_n$ and let $\mathbf{p}_1, \mathbf{p}_2, \dots, \mathbf{p}_n$ be the columns of P . Then

$AP = [A\mathbf{p}_1, A\mathbf{p}_2, \dots, A\mathbf{p}_n]$ and $PD = [\lambda_1\mathbf{p}_1, \lambda_2\mathbf{p}_2, \dots, \lambda_n\mathbf{p}_n]$. Hence $A\mathbf{p}_i = \lambda_i\mathbf{p}_i$. That is, the λ_i are the eigenvalues of A and the $\mathbf{p}_i$ are the corresponding eigenvectors. Since P is invertible, the $\mathbf{p}_i$ are linearly independent.

(if) Assume that A has n linearly independent eigenvectors $\mathbf{p}_1, \mathbf{p}_2, \dots, \mathbf{p}_n$ with corresponding eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_n$. Then $A\mathbf{p}_i = \lambda_i\mathbf{p}_i$ and reversing the above steps

$$AP = [A\mathbf{p}_1, A\mathbf{p}_2, \dots, A\mathbf{p}_n] = [\lambda_1\mathbf{p}_1, \lambda_2\mathbf{p}_2, \dots, \lambda_n\mathbf{p}_n] = PD.$$

Thus, $AP = DP$. Since the $\mathbf{p}_i$ are linearly independent, P is invertible and hence $A = P^{-1}DP$. Thus A is diagonalizable. ■

It follows from the proof of the theorem that if A is diagonalizable and has eigenvalue λ with multiplicity k , then there are k linearly independent eigenvectors associated with λ .

A matrix P is *orthogonal* if it is invertible and $P^{-1} = P^T$. A matrix A is *orthogonally diagonalizable* if there exists an orthogonal matrix P such that $P^{-1}AP = D$ is diagonal. If A is orthogonally diagonalizable, then $A = PDP^T$ and $AP = PD$. Thus, the columns of P are the eigenvectors of A and the diagonal elements of D are the corresponding eigenvalues.

If P is an orthogonal matrix, then P^TAP and A are both representations of the same linear transformation with respect to different bases. To see this, note that if $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n$ is the standard basis, then a_{ij} is the component of $A\mathbf{e}_j$ along the direction $\mathbf{e}_i$, namely, $a_{ij} = \mathbf{e}_i^T A\mathbf{e}_j$. Thus, A defines a linear transformation by specifying the image under the transformation of each basis vector. Denote by $\mathbf{p}_j$ the j^{th} column of P . It is easy to see that $(P^TAP)_{ij}$ is the component of $A\mathbf{p}_j$ along the direction $\mathbf{p}_i$, namely, $(P^TAP)_{ij} = \mathbf{p}_i^T A\mathbf{p}_j$. Since P is orthogonal, the $\mathbf{p}_j$ form a basis of space and so P^TAP represents the same linear transformation as A , but in the basis $p_1, p_2, \dots, p_n$.

Another remark is in order. Check that

$$A = PDP^T = \sum_{i=1}^n d_{ii} \mathbf{p}_i \mathbf{p}_i^T.$$

Compare this with the singular value decomposition where

$$A = \sum_{i=1}^n \sigma_i \mathbf{u}_i \mathbf{v}_i^T,$$

the only difference being that $\mathbf{u}_i$ and $\mathbf{v}_i$ can be different and indeed if A is not square, they will certainly be.

11.6.2 Symmetric Matrices

For an arbitrary matrix, some of the eigenvalues may be complex. However, for a symmetric matrix with real entries, all eigenvalues are real. The number of eigenvalues of a symmetric matrix, counting multiplicities, equals the dimension of the matrix. The set of eigenvectors associated with a given eigenvalue form a vector space. For a nonsymmetric matrix, the dimension of this space may be less than the multiplicity of the eigenvalue. Thus, a nonsymmetric matrix may not be diagonalizable. However, for a symmetric matrix the eigenvectors associated with a given eigenvalue form a vector space of dimension equal to the multiplicity of the eigenvalue. Thus, all symmetric matrices are diagonalizable. The above facts for symmetric matrices are summarized in the following theorem.

Theorem 11.7 (Real Spectral Theorem) *Let A be a real symmetric matrix. Then*

1. *The eigenvalues, $\lambda_1, \lambda_2, \dots, \lambda_n$, are real, as are the components of the corresponding eigenvectors, $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$.*
2. **Spectral Decomposition** *A is orthogonally diagonalizable and indeed*

$$A = VDV^T = \sum_{i=1}^n \lambda_i \mathbf{v}_i \mathbf{v}_i^T,$$

where V is the matrix with columns $\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_n$ and D is a diagonal matrix with entries $\lambda_1, \lambda_2, \dots, \lambda_n$.

Proof: $A\mathbf{v}_i = \lambda_i \mathbf{v}_i$ and $\mathbf{v}_i^c A \mathbf{v}_i = \lambda_i \mathbf{v}_i^c \mathbf{v}_i$. Here the c superscript means conjugate transpose. Then

$$\lambda_i = \mathbf{v}_i^c A \mathbf{v}_i = (\mathbf{v}_i^c A \mathbf{v}_i)^{cc} = (\mathbf{v}_i^c A^c \mathbf{v}_i)^c = (\mathbf{v}_i^c A \mathbf{v}_i)^c = \lambda_i^c$$

and hence λ_i is real.

Since λ_i is real, a nontrivial solution to $(A - \lambda_i I) \mathbf{x} = 0$ has real components.

Let P be a real symmetric matrix such that $P\mathbf{v}_1 = \mathbf{e}_1$ where $\mathbf{e}_1 = (1, 0, 0, \dots, 0)^T$ and $P^{-1} = P^T$. We will construct such a P shortly. Since $A\mathbf{v}_1 = \lambda_1 \mathbf{v}_1$,

$$PAP^T \mathbf{e}_1 = PA\mathbf{v}_1 = \lambda_1 P\mathbf{v}_1 = \lambda_1 \mathbf{e}_1.$$

The condition $PAP^T \mathbf{e}_1 = \lambda_1 \mathbf{e}_1$ plus symmetry implies that $PAP^T = \begin{pmatrix} \lambda_1 & 0 \\ 0 & A' \end{pmatrix}$ where A' is $n-1$ by $n-1$ and symmetric. By induction, A' is orthogonally diagonalizable. Let Q be the orthogonal matrix with $QA'Q^T = D'$, a diagonal matrix. Q is $(n-1) \times (n-1)$. Augment Q to an $n \times n$ matrix by putting 1 in the $(1, 1)$ position and 0 elsewhere in the first row and column. Call the resulting matrix R . R is orthogonal too.

$$R \begin{pmatrix} \lambda_1 & 0 \\ 0 & A' \end{pmatrix} R^T = \begin{pmatrix} \lambda_1 & 0 \\ 0 & D' \end{pmatrix} \implies RPAP^T R^T = \begin{pmatrix} \lambda_1 & 0 \\ 0 & D' \end{pmatrix}.$$

Since the product of two orthogonal matrices is orthogonal, this finishes the proof of (2) except it remains to construct P . For this, take an orthonormal basis of space containing $\mathbf{v}_1$. Suppose the basis is $\{\mathbf{v}_1, \mathbf{w}_2, \mathbf{w}_3, \dots\}$ and V is the matrix with these basis vectors as its columns. Then $P = V^T$ will do. ■

Theorem 11.8 (The fundamental theorem of symmetric matrices) *A real matrix A is orthogonally diagonalizable if and only if A is symmetric.*

Proof: (if) Assume A is orthogonally diagonalizable. Then there exists P such that $D = P^{-1}AP$. Since $P^{-1} = P^T$, we get

$$A = PDP^{-1} = PDP^T$$

which implies

$$A^T = (PDP^T)^T = PDP^T = A$$

and hence A is symmetric.

(only if) Already roved. ■

Note that a nonsymmetric matrix may not be diagonalizable, it may have eigenvalues that are not real, and the number of linearly independent eigenvectors corresponding to an eigenvalue may be less than its multiplicity. For example, the matrix

$$\begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 1 & 0 & 1 \end{pmatrix}$$

has eigenvalues 2 , $\frac{1}{2} + i\frac{\sqrt{3}}{2}$, and $\frac{1}{2} - i\frac{\sqrt{3}}{2}$. The matrix $\begin{pmatrix} 1 & 2 \\ 0 & 1 \end{pmatrix}$ has characteristic equation $(1 - \lambda)^2 = 0$ and thus has eigenvalue 1 with multiplicity 2 but has only one linearly independent eigenvector associated with the eigenvalue 1 , namely $\mathbf{x} = c \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ $c \neq 0$. Neither of these situations is possible for a symmetric matrix.

11.6.3 Extremal Properties of Eigenvalues

In this section we derive a min max characterization of eigenvalues that implies that the largest eigenvalue of a symmetric matrix A has a value equal to the maximum of $\mathbf{x}^T A \mathbf{x}$ over all vectors $\mathbf{x}$ of unit length. That is, the largest eigenvalue of A equals the 2-norm of A . If A is a real symmetric matrix there exists an orthogonal matrix P that diagonalizes A . Thus

$$P^T A P = D$$

where D is a diagonal matrix with the eigenvalues of A , $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$, on its diagonal. Rather than working with A , it is easier to work with the diagonal matrix D . This will be an important technique that will simplify many proofs.

Consider maximizing $\mathbf{x}^T A \mathbf{x}$ subject to the conditions

$$1. \sum_{i=1}^n x_i^2 = 1$$

$$2. \mathbf{r}_i^T \mathbf{x} = 0, \quad 1 \leq i \leq s$$

where the $\mathbf{r}_i$ are any set of nonzero vectors. We ask over all possible sets $\{\mathbf{r}_i | 1 \leq i \leq s\}$ of s vectors, what is the minimum value assumed by this maximum.

Theorem 11.9 (Min max theorem) For a symmetric matrix A , $\min_{\mathbf{r}_1, \dots, \mathbf{r}_s} \max_{\substack{\mathbf{x} \\ \mathbf{r}_i \perp \mathbf{x}}} (\mathbf{x}^T A \mathbf{x}) = \lambda_{s+1}$ where the minimum is over all sets $\{\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_s\}$ of s nonzero vectors and the maximum is over all unit vectors $\mathbf{x}$ orthogonal to the s nonzero vectors.

Proof: A is orthogonally diagonalizable. Let P satisfy $P^T P = I$ and $P^T A P = D$, D diagonal. Let $\mathbf{y} = P^T \mathbf{x}$. Then $\mathbf{x} = P \mathbf{y}$ and

$$\mathbf{x}^T A \mathbf{x} = \mathbf{y}^T P^T A P \mathbf{y} = \mathbf{y}^T D \mathbf{y} = \sum_{i=1}^n \lambda_i y_i^2$$

Since there is a one-to-one correspondence between unit vectors $\mathbf{x}$ and $\mathbf{y}$, maximizing $\mathbf{x}^T A \mathbf{x}$ subject to $\sum x_i^2 = 1$ is equivalent to maximizing $\sum_{i=1}^n \lambda_i y_i^2$ subject to $\sum y_i^2 = 1$. Since $\lambda_1 \geq \lambda_i$, $2 \leq i \leq n$, $\mathbf{y} = (1, 0, \dots, 0)$ maximizes $\sum_{i=1}^n \lambda_i y_i^2$ at λ_1 . Then $\mathbf{x} = P \mathbf{y}$ is the first column of P and is the first eigenvector of A . Similarly λ_n is the minimum value of $\mathbf{x}^T A \mathbf{x}$ subject to the same conditions.

Now consider maximizing $\mathbf{x}^T A \mathbf{x}$ subject to the conditions

$$1. \sum x_i^2 = 1$$

$$2. \mathbf{r}_i^T \mathbf{x} = 0$$

where the $\mathbf{r}_i$ are any set of nonzero vectors. We ask over all possible choices of s vectors what is the minimum value assumed by this maximum.

$$\min_{\mathbf{r}_1, \dots, \mathbf{r}_s} \max_{\substack{\mathbf{x} \\ \mathbf{r}_i^T \mathbf{x} = 0}} \mathbf{x}^T A \mathbf{x}$$

As above, we may work with $\mathbf{y}$. The conditions are

$$1. \sum y_i^2 = 1$$

$$2. \mathbf{q}_i^T \mathbf{y} = 0 \text{ where, } \mathbf{q}_i^T = \mathbf{r}_i^T P$$

Consider any choice for the vectors $\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_s$. This gives a corresponding set of $\mathbf{q}_i$. The $\mathbf{y}_i$ therefore satisfy s linear homogeneous equations. If we add $y_{s+2} = y_{s+3} = \dots y_n = 0$ we have $n - 1$ homogeneous equations in n unknowns $y_1, \dots, y_n$. There is at least one solution that can be normalized so that $\sum y_i^2 = 1$. With this choice of y

$$\mathbf{y}^T D \mathbf{y} = \sum \lambda_i y_i^2 \geq \lambda_{s+1}$$

since coefficients greater than or equal to $s + 1$ are zero. Thus, for any choice of $\mathbf{r}_i$ there will be a $\mathbf{y}$ such that

$$\max_{\substack{\mathbf{y} \\ \mathbf{r}_i^T \mathbf{y} = 0}} (\mathbf{y}^T P^T A P \mathbf{y}) \geq \lambda_{s+1}$$

and hence

$$\min_i \max_{\substack{\mathbf{y} \\ \mathbf{r}_i^T \mathbf{y} = 0}} (\mathbf{y}^T P^T A P \mathbf{y}) \geq \lambda_{s+1}.$$

However, there is a set of s constraints for which the minimum is less than or equal to λ_{s+1} . Fix the relations to be $y_i = 0 \quad 1 \leq i \leq s$. There are s equations in n unknowns and for any $\mathbf{y}$ subject to these relations

$$\mathbf{y}^T D \mathbf{y} = \sum_{s+1}^n \lambda_i y_i^2 \leq \lambda_{s+1}.$$

Combining the two inequalities, $\min \max \mathbf{y}^T D \mathbf{y} = \lambda_{s+1}$. ■

The above theorem tells us that the maximum of $\mathbf{x}^T A \mathbf{x}$ subject to the constraint that $|\mathbf{x}|^2 = 1$ is λ_1 . Consider the problem of maximizing $\mathbf{x}^T A \mathbf{x}$ subject to the additional restriction that $\mathbf{x}$ is orthogonal to the first eigenvector. This is equivalent to maximizing $\mathbf{y}^t P^t A P \mathbf{y}$ subject to $\mathbf{y}$ being orthogonal to $(1, 0, \dots, 0)$, i.e. the first component of $\mathbf{y}$ being 0. This maximum is clearly λ_2 and occurs for $\mathbf{y} = (0, 1, 0, \dots, 0)$. The corresponding $\mathbf{x}$ is the second column of P or the second eigenvector of A .

Similarly the maximum of $\mathbf{x}^T A \mathbf{x}$ for $\mathbf{p}_1^T \mathbf{x} = \mathbf{p}_2^T \mathbf{x} = \dots \mathbf{p}_s^T \mathbf{x} = 0$ is λ_{s+1} and is obtained for $\mathbf{x} = \mathbf{p}_{s+1}$.

11.6.4 Eigenvalues of the Sum of Two Symmetric Matrices

The min max theorem is useful in proving many other results. The following theorem shows how adding a matrix B to a matrix A changes the eigenvalues of A . The theorem is useful for determining the effect of a small perturbation on the eigenvalues of A .

Theorem 11.10 *Let A and B be $n \times n$ symmetric matrices. Let $C = A + B$. Let α_i , β_i , and γ_i denote the eigenvalues of A , B , and C respectively, where $\alpha_1 \geq \alpha_2 \geq \dots \alpha_n$ and similarly for β_i, γ_i . Then $\alpha_s + \beta_1 \geq \gamma_s \geq \alpha_s + \beta_n$.*

Proof: By the min max theorem we have

$$\alpha_s = \min_{\mathbf{r}_1, \dots, \mathbf{r}_{s-1}} \max_{\substack{\mathbf{x} \\ \mathbf{r}_i \perp \mathbf{x}}} (\mathbf{x}^T A \mathbf{x}).$$

Suppose $\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_{s-1}$ attain the minimum in the expression. Then using the min max theorem on C ,

$$\begin{aligned} \gamma_s &\leq \max_{\mathbf{x} \perp \mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_{s-1}} (\mathbf{x}^T (A + B) \mathbf{x}) \\ &\leq \max_{\mathbf{x} \perp \mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_{s-1}} (\mathbf{x}^T A \mathbf{x}) + \max_{\mathbf{x} \perp \mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{r}_{s-1}} (\mathbf{x}^T B \mathbf{x}) \\ &\leq \alpha_s + \max_{\mathbf{x}} (\mathbf{x}^T B \mathbf{x}) \leq \alpha_s + \beta_1. \end{aligned}$$

Therefore, $\gamma_s \leq \alpha_s + \beta_1$.

An application of the result to $A = C + (-B)$, gives $\alpha_s \leq \gamma_s - \beta_n$. The eigenvalues of $-B$ are minus the eigenvalues of B and thus $-\beta_n$ is the largest eigenvalue. Hence $\gamma_s \geq \alpha_s + \beta_n$ and combining inequalities yields $\alpha_s + \beta_1 \geq \gamma_s \geq \alpha_s + \beta_n$. ■

Lemma 11.11 *Let A and B be $n \times n$ symmetric matrices. Let $C = A + B$. Let α_i , β_i , and γ_i denote the eigenvalues of A , B , and C respectively, where $\alpha_1 \geq \alpha_2 \geq \dots \alpha_n$ and similarly for β_i, γ_i . Then $\gamma_{r+s-1} \leq \alpha_r + \beta_s$.*

Proof: There is a set of $r-1$ relations such that over all $\mathbf{x}$ satisfying the $r-1$ relationships

$$\max(\mathbf{x}^T A \mathbf{x}) = \alpha_r.$$

And a set of $s-1$ relations such that over all $\mathbf{x}$ satisfying the $s-1$ relationships

$$\max(\mathbf{x}^T B \mathbf{x}) = \beta_s.$$

Consider $\mathbf{x}$ satisfying all these $r+s-2$ relations. For any such $\mathbf{x}$

$$\mathbf{x}^T C \mathbf{x} = \mathbf{x}^T A \mathbf{x} + \mathbf{x}^T B \mathbf{x} \leq \alpha_r + \beta_s$$

and hence over all the $\mathbf{x}$

$$\max(\mathbf{x}^T C \mathbf{x}) \leq \alpha_r + \beta_s$$

Taking the minimum over all sets of $r+s-2$ relations

$$\gamma_{r+s-1} = \min \max(\mathbf{x}^T C \mathbf{x}) \leq \alpha_r + \beta_s$$

■

11.6.5 Norms

A set of vectors $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$ is **orthogonal** if $\mathbf{x}_i^T \mathbf{x}_j = 0$ for $i \neq j$ and is **orthonormal** if in addition $|\mathbf{x}_i| = 1$ for all i . A matrix A is **orthonormal** if $A^T A = I$. If A is a square orthonormal matrix, then rows as well as columns are orthogonal. In other words, if A is square orthonormal, then A^T is also. In the case of matrices over the complexes, the concept of an orthonormal matrix is replaced by that of a unitary matrix. A^* is the conjugate transpose of A if $a_{ij}^* = \bar{a}_{ji}$ where a_{ij}^* is the ij^{th} entry of A^* and $\bar{a}^*$ is the complex conjugate of the ij^{th} element of A . A matrix A over the field of complex numbers is **unitary** if $AA^* = I$.

Norms

A **norm** on $\mathbf{R}^n$ is a function $f : \mathbf{R}^n \rightarrow \mathbf{R}$ satisfying the following three axioms:

1. $f(\mathbf{x}) \geq 0$,
2. $f(\mathbf{x} + \mathbf{y}) \leq f(\mathbf{x}) + f(\mathbf{y})$, and
3. $f(\alpha \mathbf{x}) = |\alpha|f(\mathbf{x})$.

A norm on a vector space provides a distance function where

$$\text{distance}(\mathbf{x}, \mathbf{y}) = \text{norm}(\mathbf{x} - \mathbf{y}).$$

An important class of norms for vectors is the p -norms defined for $p > 0$ by

$$|\mathbf{x}|_p = (|\mathbf{x}_1|^p + \dots + |\mathbf{x}_n|^p)^{\frac{1}{p}}.$$

Important special cases are

$$\begin{aligned} |\mathbf{x}|_0 & \text{ the number of non zero entries} \\ |\mathbf{x}|_1 & = |x_1| + \dots + |x_n| \\ |\mathbf{x}|_2 & = \sqrt{|x_1|^2 + \dots + |x_n|^2} \\ |\mathbf{x}|_\infty & = \max |x_i|. \end{aligned}$$

Lemma 11.12 For any $1 \leq p < q$, $|\mathbf{x}|_q \leq |\mathbf{x}|_p$.

Proof:

$$|\mathbf{x}|_q^q = \sum_i |x_i|^q.$$

Let $a_i = |x_i|^q$ and $\rho = p/q$. Using the inequality (see 11.3) that for any non-negative reals $a_1, a_2, \dots, a_n$ and any $\rho \in (0, 1)$, we have $(\sum_{i=1}^n a_i)^\rho \leq \sum_{i=1}^n a_i^\rho$, the lemma is proved. ■

There are two important matrix norms, the matrix p -norm

$$\|A\|_p = \max_{|\mathbf{x}|=1} \|A\mathbf{x}\|_p$$

and the Frobenius norm

$$\|A\|_F = \sqrt{\sum_{ij} a_{ij}^2}.$$

Let $\mathbf{a}_i$ be the i^{th} column of A . Then $\|A\|_F^2 = \sum_i \mathbf{a}_i^T \mathbf{a}_i = \text{tr}(A^T A)$. A similar argument on the rows yields $\|A\|_F^2 = \text{tr}(AA^T)$. Thus, $\|A\|_F^2 = \text{tr}(A^T A) = \text{tr}(AA^T)$. If A is symmetric and rank k

$$\|A\|_2^2 \leq \|A\|_F^2 \leq k \|A\|_2^2.$$

11.6.6 Important Norms and Their Properties

Lemma 11.13 $\|AB\|_2 \leq \|A\|_2 \|B\|_2$

Proof: $\|AB\|_2 = \max_{|\mathbf{x}|=1} |AB\mathbf{x}|$. Let $\mathbf{y}$ be the value of $\mathbf{x}$ that achieves the maximum and let $\mathbf{z} = B\mathbf{y}$. Then

$$\|AB\|_2 = |AB\mathbf{y}| = |A\mathbf{z}| = \left| A \frac{\mathbf{z}}{|\mathbf{z}|} \right| |\mathbf{z}|$$

But $\left| A \frac{\mathbf{z}}{|\mathbf{z}|} \right| \leq \max_{|\mathbf{x}|=1} |A\mathbf{x}| = \|A\|_2$ and $|\mathbf{z}| \leq \max_{|\mathbf{x}|=1} |B\mathbf{x}| = \|B\|_2$. Thus $\|AB\|_2 \leq \|A\|_2 \|B\|_2$. ■

Let Q be an orthonormal matrix.

Lemma 11.14 For all $\mathbf{x}$, $|Q\mathbf{x}| = |\mathbf{x}|$.

Proof: $|Q\mathbf{x}|_2^2 = \mathbf{x}^T Q^T Q \mathbf{x} = \mathbf{x}^T \mathbf{x} = |\mathbf{x}|_2^2$. ■

Lemma 11.15 $\|QA\|_2 = \|A\|_2$

Proof: For all $\mathbf{x}$, $|Q\mathbf{x}| = |\mathbf{x}|$. Replacing $\mathbf{x}$ by $A\mathbf{x}$, $|QA\mathbf{x}| = |A\mathbf{x}|$ and thus $\max_{|\mathbf{x}|=1} |QA\mathbf{x}| = \max_{|\mathbf{x}|=1} |A\mathbf{x}|$ ■

Lemma 11.16 $\|AB\|_F^2 \leq \|A\|_F^2 \|B\|_F^2$

Proof: Let $\mathbf{a}_i$ be the i^{th} column of A and let $\mathbf{b}_j$ be the j^{th} column of B . By the Cauchy-Schwartz inequality $\|\mathbf{a}_i^T \mathbf{b}_j\| \leq \|\mathbf{a}_i\| \|\mathbf{b}_j\|$. Thus $\|AB\|_F^2 = \sum_i \sum_j |\mathbf{a}_i^T \mathbf{b}_j|^2 \leq \sum_i \sum_j \|\mathbf{a}_i\|^2 \|\mathbf{b}_j\|^2 = \sum_i \|\mathbf{a}_i\|^2 \sum_j \|\mathbf{b}_j\|^2 = \|A\|_F^2 \|B\|_F^2$ ■

Lemma 11.17 $\|QA\|_F = \|A\|_F$

Proof: $\|QA\|_F^2 = \text{Tr}(A^T Q^T Q A) = \text{Tr}(A^T A) = \|A\|_F^2$. ■

Lemma 11.18 For real, symmetric matrix A with eigenvalues $\lambda_1 \geq \lambda_2 \geq \dots$, $\|A\|_2^2 = \max(\lambda_1^2, \lambda_n^2)$ and $\|A\|_F^2 = \lambda_1^2 + \lambda_2^2 + \dots + \lambda_n^2$

Proof: Suppose the spectral decomposition of A is PDP^T , where, P is an orthogonal matrix and D is diagonal. We saw that $\|P^T A\|_2 = \|A\|_2$. Applying this again, $\|P^T A P\|_2 = \|A\|_2$. But, $P^T A P = D$ and clearly for a diagonal matrix D , $\|D\|_2$ is the largest absolute value diagonal entry from which the first equation follows. The proof of the second is analogous. ■

If A is real and symmetric and of rank k then $\|A\|_2^2 \leq \|A\|_F^2 \leq k \|A\|_2^2$

Theorem 11.19 $\|A\|_2^2 \leq \|A\|_F^2 \leq k \|A\|_2^2$

Proof: It is obvious for diagonal matrices that $\|D\|_2^2 \leq \|D\|_F^2 \leq k \|D\|_2^2$. Let $D = Q^T A Q$ where Q is orthonormal. The result follows immediately since for Q orthonormal, $\|QA\|_2 = \|A\|_2$ and $\|QA\|_F = \|A\|_F$. ■

Real and symmetric are necessary for some of these theorems. This condition was needed to express $\Sigma = Q^T A Q$. For example, suppose

$$A = \begin{pmatrix} 1 & 1 & & \\ & 1 & 1 & \\ & \vdots & \vdots & \ddots \\ & 1 & 1 & \end{pmatrix} \quad 0$$

$\|A\|_2 = 2$ and $\|A\|_F = \sqrt{2n}$. But A is rank 2 and $\|A\|_F > 2 \|A\|_2$ for $n > 8$.

Lemma 11.20 Let A be a symmetric matrix. Then $\|A\|_2 = \max_{|\mathbf{x}|=1} |\mathbf{x}^T A \mathbf{x}|$.

Proof: By definition, the 2-norm of A is $\|A\|_2 = \max_{|\mathbf{x}|=1} |A\mathbf{x}|$. Thus,

$$\|A\|_2 = \max_{|\mathbf{x}|=1} |A\mathbf{x}| = \max_{|\mathbf{x}|=1} \sqrt{\mathbf{x}^T A^T A \mathbf{x}} = \sqrt{\lambda_1^2} = \lambda_1 = \max_{|\mathbf{x}|=1} |\mathbf{x}^T A \mathbf{x}|$$

The two norm of a matrix A is greater than or equal to the 2-norm of any of its columns. Let A_u be a column of A .

Lemma 11.21 $|A_u| \leq \|A\|_2$

Proof: Let e_u be the unit vector with a 1 in position u and all other entries zero. Note $\lambda = \max_{|x|=1} |Ax|$. Let $x = e_u$ where a_u is row u . Then $|a_u| = |Ae_u| \leq \max_{|x|=1} |Ax| = \lambda$ ■

***** The following is now covered already under Singular Value Decomposition. So commented out from here *****

11.6.7 Linear Algebra

Lemma 11.22 *Let A be an $n \times n$ symmetric matrix. Then $\det(A) = \lambda_1 \lambda_2 \cdots \lambda_n$*

Proof: The $\det(A - \lambda I)$ is a polynomial in λ of degree n . The coefficient of λ^n will be ± 1 depending on whether n is odd or even. Let the roots of this polynomial be $\lambda_1, \lambda_2, \dots, \lambda_n$. Then $\det(A - \lambda I) = (-1)^n \prod_{i=1}^n (\lambda - \lambda_i)$. Thus

$$\det(A) = \det(A - \lambda I)|_{\lambda=0} = (-1)^n \prod_{i=1}^n (\lambda - \lambda_i) \Big|_{\lambda=0} = \lambda_1 \lambda_2 \cdots \lambda_n$$

■

The trace of a matrix is defined to be the sum of its diagonal elements. That is, $\text{tr}(A) = a_{11} + a_{22} + \cdots + a_{nn}$.

Lemma 11.23 $\text{tr}(A) = \lambda_1 + \lambda_2 + \cdots + \lambda_n$

Proof: Consider the coefficient of λ^{n-1} in $\det(A - \lambda I) = (-1)^n \prod_{i=1}^n (\lambda - \lambda_i)$. Write

$$A - \lambda I = \begin{pmatrix} a_{11} - \lambda & a_{12} & \cdots \\ a_{21} & a_{22} - \lambda & \cdots \\ \vdots & \vdots & \ddots \end{pmatrix}.$$

Calculate $\det(A - \lambda I)$ by expanding along the first row. Each term in the expansion involves a determinant of size $n - 1$ which is a polynomial in λ of deg $n - 2$ except for the principal minor which is of deg $n - 1$. Thus the term of deg $n - 1$ comes from

$$(a_{11} - \lambda)(a_{22} - \lambda) \cdots (a_{nn} - \lambda)$$

and has coefficient $(-1)^{n-1}(a_{11} + a_{22} + \cdots + a_{nn})$. Now

$$\begin{aligned} (-1)^n \prod_{i=1}^n (\lambda - \lambda_i) &= (-1)^n (\lambda - \lambda_1)(\lambda - \lambda_2) \cdots (\lambda - \lambda_n) \\ &= (-1)^n (\lambda^n - (\lambda_1 + \lambda_2 + \cdots + \lambda_n)\lambda^{n-1} + \cdots) \end{aligned}$$

Therefore equating coefficients $\lambda_1 + \lambda_2 + \cdots + \lambda_n = a_{11} + a_{22} + \cdots + a_{nn} = \text{tr}(A)$

Note that $(\text{tr}(A))^2 \neq \text{tr}(A^2)$. For example $A = \begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix}$ has $\text{tr} 3$, $A^2 = \begin{pmatrix} 1 & 0 \\ 0 & 4 \end{pmatrix}$ has trace $5 \neq 9$. However $\text{tr}(A^2) = \lambda_1^2 + \lambda_2^2 + \cdots + \lambda_n^2$. To see this, observe that $A^2 = (V^T D V)^2 = V^T D^2 V$. Thus, the eigenvalues of A^2 are the squares of the eigenvalues for A .

■

Alternative proof that $\text{tr}(A) = \lambda_1 + \lambda_2 + \cdots + \lambda_n$.

Suppose the spectral decomposition of A is $A = PDP^T$. We have

$$\text{Tr}(A) = \text{tr}(PDP^T) = \text{Tr}(DP^TP) = \text{Tr}(D) = \lambda_1 + \lambda_2 + \cdots + \lambda_n.$$

Lemma 11.24 *If A is $n \times m$ and B is a $m \times n$ matrix, then $\text{tr}(AB) = \text{tr}(BA)$.*

$$\text{tr}(AB) = \sum_{i=1}^n \sum_{j=1}^n a_{ij} b_{ji} = \sum_{j=1}^n \sum_{i=1}^n b_{ji} a_{ij} = \text{tr}(BA)$$

Pseudo inverse

Let A be an $n \times m$ rank r matrix and let $A = U\Sigma V^T$ be the singular value decomposition of A . Let $\Sigma' = \text{diag}\left(\frac{1}{\sigma_1}, \dots, \frac{1}{\sigma_r}, 0, \dots, 0\right)$ where $\sigma_1, \dots, \sigma_r$ are the nonzero singular values of A . Then $A' = V\Sigma'U^T$ is the pseudo inverse of A . It is the unique X that minimizes $\|AX - I\|_F$.

Second eigenvector – similar material in 4.4 p5

Suppose the eigenvalues of a matrix are $\lambda_1 \geq \lambda_2 \geq \cdots$. The second eigenvalue, namely, λ_2 plays an important role for matrices representing graphs. It may be the case that $|\lambda_n| > |\lambda_2|$.

Why is the second eigenvalue so important? Consider partitioning the vertices of a regular degree d graph $G = (V, E)$ into two blocks of equal size so as to minimize the number of edges between the two blocks. Assign value $+1$ to the vertices in one block and -1 to the vertices in the other block. Let $\mathbf{x}$ be the vector whose components are the ± 1 values assigned to the vertices. If two vertices, i and j , are in the same block, then x_i and x_j are both $+1$ or both -1 and $(x_i - x_j)^2 = 0$. If vertices i and j are in different blocks then $(x_i - x_j)^2 = 4$. Thus, partitioning the vertices into two blocks so as to minimize the edges between vertices in different blocks is equivalent to finding a vector $\mathbf{x}$ with coordinates ± 1 of which half of its coordinates are $+1$ and half of which are -1 that minimizes

$$E_{\text{cut}} = \frac{1}{4} \sum_{ij \in E} (x_i - x_j)^2$$

Let A be the adjacency matrix of G . Then

$$\begin{aligned} \mathbf{x}^T A \mathbf{x} &= \sum_{ij} a_{ij} x_i x_j = 2 \sum_{\text{edges}} x_i x_j \\ &= 2 \times \left(\begin{array}{c} \text{number of edges} \\ \text{within components} \end{array} \right) - 2 \times \left(\begin{array}{c} \text{number of edges} \\ \text{between components} \end{array} \right) \\ &= 2 \times \left(\begin{array}{c} \text{total number} \\ \text{of edges} \end{array} \right) - 4 \times \left(\begin{array}{c} \text{number of edges} \\ \text{between components} \end{array} \right) \end{aligned}$$

Maximizing $\mathbf{x}^T A \mathbf{x}$ over all $\mathbf{x}$ whose coordinates are ± 1 and half of whose coordinates are $+1$ is equivalent to minimizing the number of edges between components.

Since finding such an $\mathbf{x}$ is computational difficult, replace the integer condition on the components of $\mathbf{x}$ and the condition that half of the components are positive and half of the components are negative with the conditions $\sum_{i=1}^n x_i^2 = 1$ and $\sum_{i=1}^n x_i = 0$. Then finding the optimal $\mathbf{x}$ gives us the second eigenvalue

$$\lambda_2 = \max_{\mathbf{x} \perp \mathbf{v}_1} \frac{\mathbf{x}^T A \mathbf{x}}{\sum x_i^2}$$

Actually we should use $\sum_{i=1}^n x_i^2 = n$ not $\sum_{i=1}^n x_i^2 = 1$. Thus $n\lambda_2$ must be greater than $2 \times \left(\begin{array}{c} \text{total number} \\ \text{of edges} \end{array} \right) - 4 \times \left(\begin{array}{c} \text{number of edges} \\ \text{between components} \end{array} \right)$ since the maximum is taken over a larger set of $\mathbf{x}$. The fact that λ_2 gives us a bound on the minimum number of cross edges is what makes it so important.

11.6.8 Distance between subspaces

Define the square of the distance between two subspaces by

$$\text{dist}^2(X_1, X_2) = \|X_1 - X_2 X_2^T X_1\|_F^2$$

Since $X_1 - X_2 X_2^T X_1$ and $X_2 X_2^T X_1$ are orthogonal

$$\|X_1\|_F^2 = \|X_1 - X_2 X_2^T X_1\|_F^2 + \|X_2 X_2^T X_1\|_F^2$$

and hence

$$\text{dist}^2(X_1, X_2) = \|X_1\|_F^2 - \|X_2 X_2^T X_1\|_F^2.$$

Intuitively, the distance between X_1 and X_2 is the Frobenius norm of the component of X_1 not in the space spanned by the columns of X_2 .

If X_1 and X_2 are 1-dimensional unit length vectors, $\text{dist}^2(X_1, X_2)$ is the sin squared of the angle between the spaces.

Example: Consider two subspaces in four dimensions

$$X_1 = \begin{pmatrix} \frac{1}{\sqrt{2}} & 0 \\ 0 & \frac{1}{\sqrt{3}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{3}} \\ 0 & \frac{1}{\sqrt{3}} \end{pmatrix} \quad X_2 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \\ 0 & 0 \\ 0 & 0 \end{pmatrix}$$

Here

$$\begin{aligned} \text{dist}^2(X_1, X_2) &= \left\| \begin{pmatrix} \frac{1}{\sqrt{2}} & 0 \\ 0 & \frac{1}{\sqrt{3}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{3}} \\ 0 & \frac{1}{\sqrt{3}} \end{pmatrix} - \begin{pmatrix} 1 & 0 \\ 0 & 1 \\ 0 & 0 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \end{pmatrix} \begin{pmatrix} \frac{1}{\sqrt{2}} & 0 \\ 0 & \frac{1}{\sqrt{3}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{3}} \\ 0 & \frac{1}{\sqrt{3}} \end{pmatrix} \right\|_F^2 \\ &= \left\| \begin{pmatrix} 0 & 0 \\ 0 & 0 \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{3}} \\ 0 & \frac{1}{\sqrt{3}} \end{pmatrix} \right\|_F^2 = \frac{7}{6} \end{aligned}$$

In essence, we projected each column vector of X_1 onto X_2 and computed the Frobenius norm of X_1 minus the projection. The Frobenius norm of each column is the sin squared of the angle between the original column of X_1 and the space spanned by the columns of X_2 . ■

11.7 Generating Functions

A sequence $a_0, a_1, \dots$, can be represented by a generating function $g(x) = \sum_{i=0}^{\infty} a_i x^i$. The advantage of the generating function is that it captures the entire sequence in a closed form that can be manipulated as an entity. For example, if $g(x)$ is the generating function for the sequence $a_0, a_1, \dots$, then $x \frac{d}{dx} g(x)$ is the generating function for the sequence $0, a_1, 2a_2, 3a_3, \dots$ and $x^2 g''(x) + xg'(x)$ is the generating function for the sequence $0, a_1, 4a_2, 9a_3, \dots$

Example: The generating function for the sequence $1, 1, \dots$ is $\sum_{i=0}^{\infty} x^i = \frac{1}{1-x}$. The generating function for the sequence $0, 1, 2, 3, \dots$ is

$$\sum_{i=0}^{\infty} i x^i = \sum_{i=0}^{\infty} x \frac{d}{dx} x^i = x \frac{d}{dx} \sum_{i=0}^{\infty} x^i = x \frac{d}{dx} \frac{1}{1-x} = \frac{x}{(1-x)^2}.$$

Example: If A can be selected 0 or 1 times and B can be selected 0, 1, or 2 times and C can be selected 0, 1, 2, or 3 times, in how many ways can we select five objects. Consider the generating function for the number of ways to select objects. The generating function for the number of ways of selecting objects, selecting only A's is $1+x$, only B's is $1+x+x^2$, and only C's is $1+x+x^2+x^3$. The generating function when selecting A's, B's, and C's is the product. ■

$$(1+x)(1+x+x^2)(1+x+x^2+x^3) = 1 + 3x + 5x^2 + 6x^3 + 5x^4 + 3x^5 + x^6$$

The coefficient of x^5 is 3 and hence we can select five objects in three ways: ABBCC, ABCCC, or BBCCC. ■

The generating functions for the sum of random variables

Let $f(x) = \sum_{i=0}^{\infty} p_i x^i$ be the generating function for an integer valued random variable where p_i is the probability that the random variable takes on value i . Let $g(x) = \sum_{i=0}^{\infty} q_i x^i$ be the generating function of an independent integer valued random variable where q_i is the probability that the random variable takes on the value i . The sum of these two random variables has the generating function $f(x)g(x)$. This is because the coefficient of x^i in the product $f(x)g(x)$ is $\sum_{k=0}^i p_k q_{k-i}$ and this is also the probability that the sum of the random variables is i . Repeating this, the generating function of a sum of independent nonnegative integer valued random variables is the product of their generating functions.

11.7.1 Generating Functions for Sequences Defined by Recurrence Relationships

Consider the Fibonacci sequence

$$0, 1, 1, 2, 3, 5, 8, 13, 21, 34, 55, 89, \dots$$

defined by the recurrence relationship

$$f_0 = 0 \quad f_1 = 1 \quad f_i = f_{i-1} + f_{i-2} \quad i \geq 2$$

Multiply each side of the recurrence by x^i and sum from i equals two to infinity.

$$\begin{aligned} \sum_{i=2}^{\infty} f_i x^i &= \sum_{i=2}^{\infty} f_{i-1} x^i + \sum_{i=2}^{\infty} f_{i-2} x^i \\ f_2 x^2 + f_3 x^3 + \dots &= f_1 x^2 + f_2 x^3 + \dots + f_0 x^2 + f_1 x^3 + \dots \\ f_2 x^2 + f_3 x^3 + \dots &= x(f_1 x + f_2 x^2 + \dots) + x^2(f_0 + f_1 x + \dots) \end{aligned} \quad (11.1)$$

Let

$$f(x) = \sum_{i=0}^{\infty} f_i x^i. \quad (11.2)$$

Substituting (11.2) into (11.1) yields

$$\begin{aligned} f(x) - f_0 - f_1 x &= x(f(x) - f_0) + x^2 f(x) \\ f(x) - x &= x f(x) + x^2 f(x) \\ f(x)(1 - x - x^2) &= x \end{aligned}$$

Thus, $f(x) = \frac{x}{1-x-x^2}$ is the generating function for the Fibonacci sequence.

Note that generating functions are formal manipulations and do not necessarily converge outside some region of convergence. Consider the generating function $f(x) = \sum_{i=0}^{\infty} f_i x^i = \frac{x}{1-x-x^2}$ for the Fibonacci sequence. Using $\sum_{i=0}^{\infty} f_i x^i$,

$$f(1) = f_0 + f_1 + f_2 + \cdots = \infty$$

and using $f(x) = \frac{x}{1-x-x^2}$

$$f(1) = \frac{1}{1-1-1} = -1.$$

Asymptotic behavior

To determine the asymptotic behavior of the Fibonacci sequence write

$$f(x) = \frac{x}{1-x-x^2} = \frac{\frac{\sqrt{5}}{5}}{1-\phi_1 x} + \frac{-\frac{\sqrt{5}}{5}}{1-\phi_2 x}$$

where $\phi_1 = \frac{1+\sqrt{5}}{2}$ and $\phi_2 = \frac{1-\sqrt{5}}{2}$ are the two roots of the quadratic $1-x-x^2=0$.

Then

$$f(x) = \frac{\sqrt{5}}{5} (1 + \phi_1 x + (\phi_1 x)^2 + \cdots - (1 + \phi_2 x + (\phi_2 x)^2 + \cdots)).$$

Thus,

$$f_n = \frac{\sqrt{5}}{5} (\phi_1^n - \phi_2^n).$$

Since $\phi_2 < 1$ and $\phi_1 > 1$, for large n , $f_n \cong \frac{\sqrt{5}}{5} \phi_1^n$. In fact, since $f_n = \frac{\sqrt{5}}{5} (\phi_1^n - \phi_2^n)$ is an integer and $\phi_2 < 1$, it must be the case that $f_n = \left\lfloor f_n + \frac{\sqrt{5}}{5} \phi_2^n \right\rfloor$. Hence $f_n = \left\lfloor \frac{\sqrt{5}}{5} \phi_1^n \right\rfloor$ for all n .

Means and standard deviations of sequences

Generating functions are useful for calculating the mean and standard deviation of a sequence. Let z be an integral valued random variable where p_i is the probability that z equals i . The expected value of z is given by $m = \sum_{i=0}^{\infty} i p_i$. Let $p(x) = \sum_{i=0}^{\infty} p_i x^i$ be the generating function for the sequence $p_1, p_2, \dots$. The generating function for the sequence $p_1, 2p_2, 3p_3, \dots$ is

$$x \frac{d}{dx} p(x) = \sum_{i=0}^{\infty} i p_i x^i.$$

Thus, the expected value of the random variable z is $m = x p'(x)|_{x=1} = p'(1)$. If p was not a probability function, its average value would be $\frac{p'(1)}{p(1)}$ since we would need to normalize

the area under p to one.

The second moment of z , is $E(z^2) - E^2(z)$ and can be obtained as follows.

$$\begin{aligned} x^2 \frac{d}{dx} p(x) \Big|_{x=1} &= \sum_{i=0}^{\infty} i(i-1)x^i p(x) \Big|_{x=1} \\ &= \sum_{i=0}^{\infty} i^2 x^i p(x) \Big|_{x=1} - \sum_{i=0}^{\infty} i x^i p(x) \Big|_{x=1} \\ &= E(z^2) - E(z). \end{aligned}$$

Thus, $\sigma^2 = p''(1) + p'(1) - (p'(1))^2$.

11.7.2 The Exponential Generating Function and the Moment Generating Function

Besides the ordinary generating function there are a number of other types of generating functions. One of these is the exponential generating function. Given a sequence $a_0, a_1, \dots$, the *associated exponential generating function* is $g(x) = \sum_{i=0}^{\infty} a_i \frac{x^i}{i!}$.

Moment generating functions

The k^{th} moment of a random variable x around the point b is given by $E((x-b)^k)$. Usually the word moment is used to denote the moment around the value 0 or around the mean. In the following, we use moment to mean the moment about the origin.

The *moment generating function* of a random variable x is defined by

$$\Psi(t) = E(e^{tx}) = \int_{-\infty}^{\infty} e^{tx} p(x) dx$$

Replacing e^{tx} by its power series expansion $1 + tx + \frac{(tx)^2}{2!} + \dots$ gives

$$\Psi(t) = \int_{-\infty}^{\infty} \left(1 + tx + \frac{(tx)^2}{2!} + \dots \right) p(x) dx$$

Thus, the k^{th} moment of x about the origin is $k!$ times the coefficient of t^k in the power series expansion of the moment generating function. Hence, the moment generating function is the exponential generating function for the sequence of moments about the origin.

The moment generating function transforms the probability distribution $p(x)$ into a function $\Psi(t)$ of t . Note $\Psi(0) = 1$ and is the area or integral of $p(x)$. The moment

generating function is closely related to the *characteristic function* which is obtained by replacing e^{tx} by e^{itx} in the above integral where $i = \sqrt{-1}$ and is related to the *Fourier transform* which is obtained by replacing e^{tx} by e^{-itx} .

$\Psi(t)$ is closely related to the Fourier transform and its properties are essentially the same. In particular, $p(x)$ can be uniquely recovered by an inverse transform from $\Psi(t)$. More specifically, if all the moments m_i are finite and the sum $\sum_{i=0}^{\infty} \frac{m_i}{i!} t^i$ converges absolutely in a region around the origin, then $p(x)$ is uniquely determined. See exercise [2].

The Gaussian probability distribution with zero mean and unit variance is given by $p(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}$. Its moments are given by

$$u_n = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} x^n e^{-\frac{x^2}{2}} dx$$

$$= \begin{cases} \frac{n!}{2^{\frac{n}{2}} (\frac{n}{2})!} & \text{n even} \\ 0 & \text{n odd} \end{cases}$$

To derive the above, use integration by parts to get $u_n = (n-1) u_{n-2}$ and combine this with $u_0 = 1$ and $u_1 = 0$. The steps are as follows. Let $u = e^{-\frac{x^2}{2}}$ and $v = x^{n-1}$. Then $u' = -xe^{-\frac{x^2}{2}}$ and $v' = (n-1)x^{n-2}$. Now $uv = \int u'v + \int uv'$ or

$$e^{-\frac{x^2}{2}} x^{n-1} = \int x^n e^{-\frac{x^2}{2}} dx + \int (n-1) x^{n-2} e^{-\frac{x^2}{2}} dx$$

From which

$$\int_{-\infty}^{\infty} x^n e^{-\frac{x^2}{2}} dx = (n-1) \int_{-\infty}^{\infty} x^{n-2} e^{-\frac{x^2}{2}} dx - e^{-\frac{x^2}{2}} x^{n-1}$$

$$\int_{-\infty}^{\infty} x^n e^{-\frac{x^2}{2}} dx = (n-1) \int_{-\infty}^{\infty} x^{n-2} e^{-\frac{x^2}{2}} dx$$

Thus, $u_n = (n-1) u_{n-2}$.

The moment generating function is given by

$$g(s) = \sum_{n=0}^{\infty} \frac{u_n s^n}{n!} = \sum_{\substack{n=0 \\ n \text{ even}}}^{\infty} \frac{n!}{2^{\frac{n}{2}} \frac{n!}{2}} \frac{s^n}{n!} = \sum_{i=0}^{\infty} \frac{s^{2i}}{2^i i!} = \sum_{i=0}^{\infty} \frac{1}{i!} \left(\frac{s^2}{2} \right)^i = e^{\frac{s^2}{2}}$$

For the general Gaussian, the moment generating function is

$$g(s) = e^{su + \left(\frac{\sigma^2}{2}\right)s^2}$$

Thus, given two independent Gaussians with mean u_1 and u_2 and variances σ_1^2 and σ_2^2 , the product of their moment generating functions is

$$e^{s(u_1+u_2)+(\sigma_1^2+\sigma_2^2)s^2},$$

the moment generating function for a Gaussian with mean $u_1 + u_2$ and variance $\sigma_1^2 + \sigma_2^2$. Thus, the convolution of two Gaussians is a Gaussian and the sum of two random variables that are both Gaussian is a Gaussian random variable.

11.8 Miscellaneous

Wigner semicircular law

If A is a symmetric, random $n \times n$ matrix with independent, mean zero, variance σ^2 entries whose absolute values are bounded by one, then with high probability the eigenvalues will be bounded by $\pm 2\sqrt{n}\sigma$

TD-IDF

The term TD-IDF stands for (term frequency) $\times$ (inverse document frequency) and is a technique for weighting keywords from a document. Let

w_{ik} be the weight of keyword k in document i ,

n_k be the number of documents containing keyword k ,

tf_{ik} be the number of occurrences of keyword k in document i , and

ndoc be the total number of documents.

Then $w_{ik} = tf_{ik} \log \frac{\text{ndoc}}{n_k}$. Sometimes the term frequency tf is normalized by the number of words in the document to prevent a bias towards long documents.

PCA and LSI

PCA applies the SVD to $X - E(X)$ whereas LSI applies the SVD to X itself.

Maximization

Maximize $f = a + 2b$ subject to $a^2 + b^2 = 1$.

$$f = a + 2b = a + 2\sqrt{1 - a^2}$$

$$\frac{\partial f}{\partial a} = 1 - (1 - a^2)^{\frac{1}{2}} 2a = 0$$

$$\frac{2a}{\sqrt{1-a^2}} = 1$$

$$4a^2 = 1 - a^2$$

$$a = \frac{1}{\sqrt{5}}$$

Lagrangian multipliers

Lagrangian multipliers are used to convert a constrained optimization problem into an unconstrained optimization. Suppose we wished to maximize a function $f(x, y)$ subject to a constraint $g(x, y) = c$. The value of $f(x, y)$ along the constraint $g(x, y) = c$ might increase for a while and then start to decrease. At the point where $f(x, y)$ stops increasing and starts to decrease, the contour line for $f(x, y)$ is tangent to the curve of the constraint $g(x, y) = c$. Stated another way the gradient of $f(x, y)$ and the gradient of $g(x, y)$ are parallel.

By introducing a new variable λ we can express the condition by $\nabla_{xy}f = \lambda \nabla_{xy}g$ and $g=c$. These two conditions hold if and only if

$$\nabla_{xy\lambda} (f(x, y) + \lambda(g(x, y) - c)) = 0$$

The partial with respect to λ establishes that $g(x, y) = c$. We have converted the constrained optimization problem in variables x and y to an unconstrained problem with variables x , y , and λ .

11.8.1 Finite Fields

For a prime p and integer n there is a unique finite field with p^n elements. In Section 7.1.4 we used the field $\text{GF}(2^n)$, which consists of polynomials of degree less than or equal to n with coefficients over the field $\text{GF}(2)$. In $\text{GF}(2^8)$

$$(x^7 + x^5 + x) + (x^6 + x^5 + x^4) = x^7 + x^6 = x^4 = x$$

Multiplication is modulo an irreducible polynomial. Thus

$$\begin{aligned} (x^7 + x^5 + x)(x^6 + x^5 + x^4) &= x^{13} + x^{12} + x^{11} + x^{11} + x^{10} + x^9 + x^7 + x^6 + x^5 \\ &= x^{13} + x^{12} + x^{10} + x^9 + x^7 + x^6 + x^5 \\ &= x^6 + x^4 + x^3 + x^2 \quad \text{mod } x^8 + x^4 + x^3 + x + 1 \end{aligned}$$

$$\begin{array}{rcccccccccccc}
-x^5(x^8 + x^4 + x^3 + x^2 + 1) = & x^{13} & +x^{12} & +x^{10} & +x^9 & & +x^7 & +x^6 & +x^5 \\
& x^{13} & & & +x^9 & +x^8 & & +x^6 & +x^5 \\
\hline
-x^4(x^8 + x^4 + x^3 + x^2 + 1) = & & x^{12} & +x^{10} & & +x^8 & +x^7 & & & & \\
& & x^{12} & & & +x^8 & +x^7 & & +x^5 & +x^4 \\
\hline
-x^2(x^8 + x^4 + x^3 + x^2 + 1) = & & & x^{10} & & & & & +x^5 & x^4 \\
& & & x^{10} & & & & x^6 & +x^5 & & x^3 & x^2 \\
\hline
& & & & & & & x^6 & & +x^4 & +x^3 & +x^2
\end{array}$$

11.8.2 Variational Methods

Expressing the log function as a minimization problem

The function $\ln(x)$ is concave and thus for a line λx with slope λ there is an intercept c such that $\lambda x + c$ is tangent to the curve $\ln(x)$ in exactly one point. The value of c is $-\ln \lambda - 1$. To see this, note that the slope of $\ln(x)$ is $1/x$ and this equals the slope of the line at $x = \frac{1}{\lambda}$. The value of c for which $\lambda x + c = \ln(x)$ at $x = 1/2$ is $c = \ln\left(\frac{1}{\lambda}\right) - 1 = -\ln \lambda - 1$.

Now for every λ and every x , $\ln(x) \leq (\lambda x - \ln \lambda - 1)$ and for every x there is a λ such that equality holds. Thus,

$$\ln(x) = \min_{\lambda} (\lambda x - \ln \lambda - 1)$$

The logistic function $f(x) = \frac{1}{1+e^{-x}}$ is not convex but we can transform it to a domain where the function is concave by taking the $\ln$ of $f(x)$.

$$g(x) = -\ln(1 + e^{-x})$$

Thus, the log logistic function can be bounded with linear functions

$$g(x) = \min_{\lambda} \{\lambda x - H(\lambda)\}$$

Taking exponential of both sides

$$f(x) = \min_{\lambda} \{e^{\lambda x - H(\lambda)}\}.$$

Convex duality

A convex function $f(x)$ can be represented via a conjugate or dual function as follows

$$f(x) = \min_{\lambda} (\lambda x - f^*(\lambda))$$

The conjugate function $f^*(x)$ is obtained from the dual expression

$$f^*(\lambda) = \min_x (\lambda x - f(x))$$

11.8.3 Hash Functions

Universal Hash Families

ADD PARAGRAPH ON MOTIVATION integrate material with Chapter

Let $M = \{1, 2, \dots, m\}$ and $N = \{1, 2, \dots, n\}$ where $m \geq n$. A family of hash functions $H = \{h | h : M \rightarrow N\}$ is said to be 2-universal if for all x and y , $x \neq y$, and for h chosen uniformly at random from H ,

$$\text{Prob}[h(x) = h(y)] \leq \frac{1}{n}$$

Note that if H is the set of all possible mappings from M to N , then H is 2-universal. In fact $\text{Prob}[h(x) = h(y)] = \frac{1}{n}$. The difficulty in letting H consist of all possible functions is that a random h from H has no short representation. What we want is a small set H where each $h \in H$ has a short representation and is easy to compute.

Note that for a 2-universal H , for any two elements x and y , $h(x)$ and $h(y)$ behave as independent random variables. For a random f and any set X the set $\{f(x) | x \in X\}$ is a set of independent random variables.

11.8.4 Application of Mean Value Theorem

The mean value theorem states that if $f(x)$ is continuous and differentiable on the interval $[a, b]$, then there exists c , $a \leq c \leq b$ such that $f'(c) = \frac{f(b)-f(a)}{b-a}$. That is, at some point between a and b the derivative of f equals the slope of the line from $f(a)$ to $f(b)$. See Figure ??.

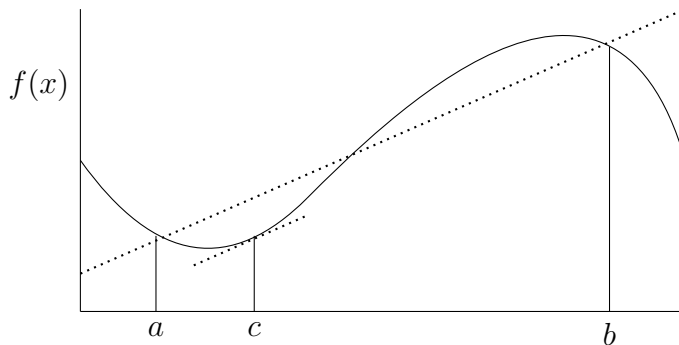


Figure 11.3: Illustration of the mean value theorem.

One application of the mean value theorem is with the Taylor expansion of a function. The Taylor expansion about the origin of $f(x)$ is

$$f(x) = f(0) + f'(0)x + \frac{1}{2!}f''(0)x^2 + \frac{1}{3!}f'''(0)x^3 + \dots \quad (11.3)$$

By the mean value theorem there exists c , $0 \leq c \leq x$, such that $f'(c) = \frac{f(x)-f(0)}{x}$ or $f(x) - f(0) = xf'(c)$. Thus

$$xf'(c) = f'(0)x + \frac{1}{2!}f''(0)x^2 + \frac{1}{3!}f'''(0)x^3 + \dots$$

and

$$f(x) = f(0) + xf'(c).$$

One could apply the mean value theorem to $f'(x)$ in

$$f'(x) = f'(0) + f''(0)x + \frac{1}{2!}f'''(0)x^2 + \dots$$

Then there exists d , $0 \leq d \leq x$ such that

$$xf''(d) = f''(0)x + \frac{1}{2!}f'''(0)x^2 + \dots$$

Integrating

$$\frac{1}{2}x^2 f''(d) = \frac{1}{2!}f''(0)x + \frac{1}{3!}f'''(0)x^3 + \dots$$

Substituting into Eq(11.3)

$$f(x) = f(0) + f'(0)x + \frac{1}{2}x^2 f''(d).$$

11.8.5 Catalan numbers

In Section?? we derived the Catalan numbers using generating functions to illustrate the use of generating functions. The formula for the Catalan numbers can be derived more easily as follows. The number of strings of length $2n$ with equal numbers of left and right parentheses is $\binom{2n}{n}$. Each of these strings is a string of balanced parentheses unless there is a prefix with one more right parentheses than left. Consider such a string and flip all parentheses following this one, from right parentheses to left and vice versa. This results in a string of $n - 1$ left parentheses and $n + 1$ right parentheses of which there are $\binom{2n}{n-1}$. Thus, there is a one-to-one correspondence between strings of length $2n$ with equal numbers of left and right parentheses that are not balanced and strings of length $2n$ with $n - 1$ left parentheses and $n + 1$ right parentheses.

We showed the correspondence from unbalanced strings of length $2n$ with equal numbers of parentheses to strings of length $2n$ with $n - 1$ left parentheses. To see the correspondence the other way take any string of length $2n$ with $n + 1$ right parentheses. Find the first position where there is one more right parentheses than left and flip all parentheses beyond this point. Clearly this results in an unbalanced string. Thus.

$$c_n = \binom{2n}{n} - \binom{2n}{n-1} = \frac{1}{n+1} \binom{2n}{n}$$

11.8.6 Sperner's Lemma

Consider a triangulation of a 2-dimensional simplex. Let the vertices of the simplex be colored R, B, and G. If the vertices on each edge of the simplex are colored only with the two colors at the endpoints then the triangulation must have a triangle whose vertices are three different colors. In fact, it must have an odd number of such vertices. A generalization of the lemma to higher dimensions also holds.

Create a graph whose vertices correspond to the triangles of the triangulation plus an additional vertex corresponding to the outside region. Connect two vertices of the graph by an edge if the triangles corresponding to the two vertices share a common edge that is color R and B. The edge of the original simplex must have an odd number of such triangular edges. Thus, the outside vertex of the graph must be of odd degree. The graph must have an even number of odd degree vertices. Each odd vertex is of degree 0, 1, or 2. The vertices of odd degree, i.e. degree one, correspond to triangles which have all three colors.

11.8.7 Prüfer

Here we prove that the number of labeled trees with n vertices is n^{n-2} . By a labeled tree we mean a tree with n vertices and n distinct labels, each label assigned to one vertex.

Theorem 11.25 *The number of labeled trees with n vertices is n^{n-2} .*

Proof: (Prüfer sequence) There is a one-to-one correspondence between labeled trees and sequences of length $n - 2$ of integers between 1 and n . An integer may repeat in the sequence. The number of such sequences is clearly n^{n-2} . Although each vertex of the tree has a unique integer label the corresponding sequence has repeating labels. The reason for this is that the labels in the sequence refer to interior vertices of the tree and the number of times the integer corresponding to an interior vertex occurs in the sequence is related to the degree of the vertex. Integers corresponding to leaves do not appear in the sequence.

To see the one-to-one correspondence, first convert a tree to a sequence by deleting the lowest numbered leaf. If the lowest numbered leaf is i and its parent is j , append j to the tail of the sequence. Repeating the process until only two vertices remain yields the

sequence. Clearly a labeled tree gives rise to only one sequence.

It remains to show how to construct a unique tree from a sequence. The proof is by induction on n . For $n = 1$ or 2 the induction hypothesis is trivially true. Assume the induction hypothesis true for $n - 1$. Certain numbers from 1 to n do not appear in the sequence and these numbers correspond to vertices that are leaves. Let i be the lowest number not appearing in the sequence and let j be the first integer in the sequence. Then i corresponds to a leaf connected to vertex j . Delete the integer j from the sequence. By the induction hypothesis there is a unique labeled tree with integer labels $1, \dots, i - 1, i + 1, \dots, n$. Add the leaf i by connecting the leaf to vertex j . We need to argue that no other sequence can give rise to the same tree. Suppose some other sequence did. Then the i^{th} integer in the sequence must be j . By the induction hypothesis the sequence with j removed is unique.

Algorithm

Create leaf list - the list of labels not appearing in the Prfer sequence. n is the length of the Prfer list plus two.

while Prfer sequence is non empty do

begin

p = first integer in Prfer sequence

e = smallest label in leaf list

Add edge (p, e)

Delete e from leaf list

Delete p from Prfer sequence

If p no longer appears in Prfer sequence add p to leaf list

end

there are two vertices e and f on leaf list, add edge (e, f)

11.9 Exercises

Exercise 11.1 What is the difference between saying $f(n)$ is $O(n^3)$ and $f(n)$ is $o(n^3)$?

Exercise 11.2 If $f(n) \sim g(n)$ what can we say about $f(n) + g(n)$ and $f(n) - g(n)$?

Exercise 11.3 What is the difference between $\sim$ and Θ ?

Exercise 11.4 If $f(n)$ is $O(g(n))$ does this imply that $g(n)$ is $\Omega(f(n))$?

Exercise 11.5 What is $\lim_{k \rightarrow \infty} \left(\frac{k-1}{k-2} \right)^{k-2}$.

Exercise 11.6 Select a , b , and c uniformly at random from $[0, 1]$. The probability that $b < a$ is $1/2$. The probability that $c < a$ is $1/2$. However, the probability that both b and c are less than a is $\frac{1}{3}$ not $1/4$. Why is this? Note that the six possible permutations abc , acb , bac , cab , bca , and cba , are all equally likely. Assume that a , b , and c are drawn from the interval $(0, 1]$. Given that $b < a$, what is the probability that $c < a$?

Exercise 11.7 Give an example of three random variables that are pairwise independent but not fully independent.

Solution: The set of variables in the following table are pairwise independent but not truly independent since 001 never occurs.

x	y	z
0	0	0
0	1	1
1	0	1
1	1	0

$$p(x=0) = \frac{1}{2} \quad p(x=0/y=0) = \frac{p(x=0, y=0)}{p(y=0)} = \frac{\frac{1}{4}}{\frac{1}{2}} = \frac{1}{2}$$

$$p(x=1) = \frac{1}{2} \quad p(x=1/y=0) = \frac{p(x=1, y=0)}{p(y=0)} = \frac{\frac{1}{4}}{\frac{1}{2}} = \frac{1}{2}$$

The other situations are similar. Thus, x and y are pairwise independent. ■

Exercise 11.8 Give examples of nonnegative valued random variables with median $>>$ mean. Can we have median $<<$ mean?

Solution:

- 1,5,6 median = 5, mean = 4, median $>$ mean
- $1, 3 \times 10^6, 3 \times 10^6 + 2$ median = 3×10^6 , mean = $2 \times 10^6 + 1$, median $>>$ mean
- 1,2,6 median = 2, mean = 3, median $<$ mean
- $1, 2, 3 \times 10^6$ median = 2, mean = $10^6 + 1$, median $<<$ mean

Exercise 11.9 Consider n samples $x_1, x_2, \dots, x_n$ from a Gaussian distribution of mean μ and variance σ . For this distribution $m = \frac{x_1+x_2+\dots+x_n}{n}$ is an unbiased estimator of μ . If μ is known then $\frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2$ is an unbiased estimator of σ^2 . Prove that if we approximate μ by m , then $\frac{1}{n-1} \sum_{i=1}^n (x_i - m)^2$ is an unbiased estimator of σ^2 .

Exercise 11.10 Given the distribution $\frac{1}{\sqrt{2\pi}3} e^{-\frac{1}{2}(\frac{x}{3})^2}$ what is the probability that $x > 1$?

Exercise 11.11 $e^{-\frac{x^2}{2}}$ has value 1 at $x = 0$ and drops off very fast as x increases. Suppose we wished to approximate $e^{-\frac{x^2}{2}}$ by a function $f(x)$ where

$$f(x) = \begin{cases} 1 & |x| \leq a \\ 0 & |x| > a \end{cases}.$$

What value of a should we use? What is the integral of the error between $f(x)$ and $e^{-\frac{x^2}{2}}$?

Exercise 11.12 Given two sets of red and black balls with the number of red and black balls in each set shown in the table below.

	red	black
Set 1	40	60
Set 2	50	50

Randomly draw a ball from one of the sets. Suppose that it turns out to be red. What is the probability that it was drawn from Set 1?

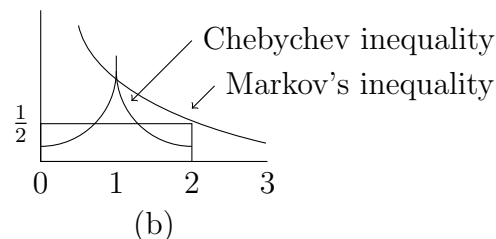
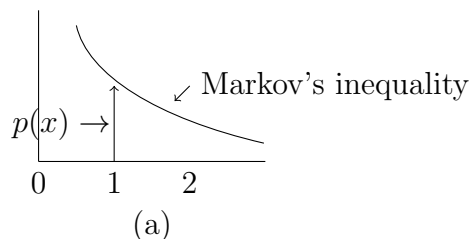
Exercise 11.13 Why cannot one prove an analogous type of theorem that states $p(x \leq a) \leq \frac{E(x)}{a}$?

Exercise 11.14 Compare the Markov and Chebyshev bounds for the following probability distributions

1. $p(x) = \begin{cases} 1 & x = 1 \\ 0 & \text{otherwise} \end{cases}$
2. $p(x) = \begin{cases} 1/2 & 0 \leq x \leq 2 \\ 0 & \text{otherwise} \end{cases}$

Solution:

1. By Markov's inequality $p(x > a) \leq \frac{1}{a}$. By Chebyshev's inequality since $\sigma^2 = 0$, $\text{Prob}(|x - 1| > 0) = 0$, hence $x = 1$.
2. By Markov's inequality $p(x > a) \leq \frac{1}{a}$.
 $\sigma^2 = E((x - 1)^2) = E(x^2 - 2x + 1)$.
 $\frac{1}{2} \int_0^2 x^2 dx = \frac{1}{6} x^3 \Big|_0^2 = \frac{4}{3}$.
 Thus $\sigma^2 = \frac{4}{3} - 2 + 1 = \frac{1}{3}$. By Chebyshev's inequality $\text{Prob}(|x - 1| > a) < \frac{1}{3a^2}$.



■

Exercise 11.15 Let s be the sum of n independent random variables $x_1, x_2, \dots, x_n$ where for each i

$$x_i = \begin{cases} 0 & \text{Prob } p \\ 1 & \text{Prob } 1-p \end{cases}$$

1. How large must δ be if we wish to have $\text{Prob}(s < (1 - \delta)m) < \varepsilon$?

2. If we wish to have $\text{Prob}(s > (1 + \delta)m) < \varepsilon$?

Exercise 11.16 What is the expected number of flips of a coin until a head is reached? Assume p is probability of a head on an individual flip. What is value if $p=1/2$?

Exercise 11.17 Given the joint probability

P(A,B)	A=0	A=1
B=0	1/16	1/8
B=1	1/4	9/16

1. What is the marginal probability of A ? of B ?

2. What is the conditional probability of B given A ?

Exercise 11.18 Consider independent random variables x_1, x_2 , and x_3 , each equal to zero with probability $\frac{1}{2}$. Let $S = x_1 + x_2 + x_3$ and let F be event that $S \in \{1, 2\}$. Conditioning on F , the variables x_1, x_2 , and x_3 are still each zero with probability $\frac{1}{2}$. Are they still independent?

Solution: No

$$\begin{aligned} \text{Prob}(x_1 = 0, x_2 = 0|F) &= 1/6 \\ \frac{1}{4} &= \text{Prob}(x_1 = 0|F)\text{Prob}(x_2 = 0|F) = \frac{1}{4} \end{aligned}$$

■

Exercise 11.19 Consider rolling two dice A and B . What is the probability that the sum S will add to nine? What is the probability that the sum will be 9 if the roll of A is 3?

Exercise 11.20 Write the generating function for the number of ways of producing chains using only pennies, nickels, and dimes. In how many ways can you produce 23 cents?

Solution: The generating functions to produce change using only pennies, nickels and dimes are

$$\begin{aligned} &(1 + x + x^2 + \cdots) \\ &(1 + x^5 + x^{10} + \cdots) \\ &(1 + x^{10} + x^{20} + \cdots) \end{aligned}$$

Thus, the generating function for the number of ways of making change using pennies, nickels, and dimes is

$$\begin{aligned} &(1 + x + x^2 + \cdots)(1 + x^5 + x^{10} + \cdots)(1 + x^{10} + x^{20} + \cdots) = \\ &1 + x + x^2 + x^3 + x^4 + 2x^5 + 2x^6 + 2x^7 + 2x^8 + 2x^9 + 4x^{10} \\ &+ 4x^{11} + 4x^{12} + 4x^{13} + 4x^{14} + 6x^{15} + 6x^{16} + 6x^{17} + 6x^{18} \\ &+ 6x^{19} + 9x^{20} + 9x^{21} + 9x^{22} + 9x^{23} + \cdots \end{aligned}$$

There are nine ways to give 23 cents change. ■

Exercise 11.21 A dice has six faces, each face of the dice having one of the numbers 1 through 6. The result of a role of the dice is the integer on the top face. Consider two roles of the dice. In how many ways can an integer be the sum of two roles of the dice.

Exercise 11.22 If $a(x)$ is the generating function for the sequence $a_0, a_1, a_2, \dots$, for what sequence is $a(x)(1-x)$ the generating function.

Exercise 11.23 How many ways can one draw n a 's and b 's with an even number of a 's.

Exercise 11.24 Find the generating function for the recurrence $a_i = 2a_{i-1} + i$ where $a_0 = 1$.

Exercise 11.25 Find a closed form for the generating function for the infinite sequence of prefect squares 1, 4, 9, 16, 25, ...

Exercise 11.26 Given that $\frac{1}{1-x}$ is the generating function for the sequence 1, 1, ..., for what sequence is $\frac{1}{1-2x}$ the generating function?

Exercise 11.27 Find a closed form for the exponential generating function for the infinite sequence of prefect squares 1, 4, 9, 16, 25, ...

Exercise 11.28 The generating function for the Catalan numbers is $c(x) = \frac{1-\sqrt{1-4x}}{2x}$. Find an approximation for c_n by expanding $\sqrt{1-4x}$ by the binomial expression and using Sterling's approximation for $n!$.

Exercise 11.29 Show that C_n is the number of ways a convex polygon with $n+2$ vertices can be partitioned into triangles by adding edges between pairs of vertices.

Exercise 11.30 Show that C_n satisfies the recurrence relationship $C_{n+1} = \frac{2(2n+1)}{n+2}C_n$

Exercise 11.31 Prove that the L_2 norm of $(a_1, a_2, \dots, a_n)$ is less than or equal to the L_1 norm of $(a_1, a_2, \dots, a_n)$.

Exercise 11.32 Prove that there exists a y , $0 \leq y \leq x$, such that $f(x) = f(0) + f'(y)x$.

Exercise 11.33 Let A be the adjacency matrix of an undirected graph G . Prove that eigenvalue λ_1 of A is at least the average degree of G .

Exercise 11.34 Show that if A is a symmetric matrix and λ_1 and λ_2 are distinct eigenvalues then their corresponding eigenvectors x_1 and x_2 are orthogonal.

Hint:

Exercise 11.35 Show that a matrix is rank k if and only if it has k nonzero eigenvalues and eigenvalue 0 of rank $n-k$.

Exercise 11.36 Prove that maximizing $\frac{x^T Ax}{x^T x}$ is equivalent to maximizing $x^T Ax$ subject to the condition that x be of unit length.

Exercise 11.37 Let A be a symmetric matrix with smallest eigenvalue $\lambda_{\min}$. Give a bound on the largest element of A^{-1} .

Exercise 11.38 Let A be the adjacency matrix of an n vertex clique with no self loops. Thus, each row of A is all ones except for the diagonal entry which is zero. What is the spectrum of A .

Exercise 11.39 Let A be the adjacency matrix of an undirect graph G . Prove that the eigenvalue λ_1 of A is at least the average degree of G .

Exercise 11.40 We are given the probability distribution for two random vectors x and y and we wish to stretch space to maximize the expected distance between them. Thus, we will multiply each coordinate by some quantity a_i . We restrict $\sum_{i=1}^d a_i^2 = d$. Thus, if we increase some coordinate by $a_i > 1$, some other coordinate must shrink. Given random vectors $x = (x_1, x_2, \dots, x_d)$ and $y = (y_1, y_2, \dots, y_d)$ how should we select a_i to maximize $E(|x - y|^2)$? The a_i stretch different coordinates. Assume

$$y_i = \begin{cases} 0 & \frac{1}{2} \\ 1 & \frac{1}{2} \end{cases}$$

and that x_i has some arbitrary distribution.

$$\begin{aligned} E(|x - y|^2) &= E \sum_{i=1}^d [a_i^2 (x_i - y_i)^2] = \sum_{i=1}^d a_i^2 E(x_i^2 - 2x_i y_i + y_i^2) \\ &= \sum_{i=1}^d a_i^2 E(x_i^2 - x_i + \frac{1}{2}) \end{aligned}$$

Since $E(x_i^2) = E(x_i)$ we get . Thus, weighting the coordinates has no effect assuming $\sum_{i=1}^d a_i^2 = 1$. Why is this? Since $E(y_i) = \frac{1}{2}$.

$E(|x - y|^2)$ is independent of the value of x_i hence its distribution.

What if $y_i = \begin{cases} 0 & \frac{3}{4} \\ 1 & \frac{1}{4} \end{cases}$ and $E(y_i) = \frac{1}{4}$. Then

$$\begin{aligned} E(|x - y|^2) &= \sum_{i=1}^d a_i^2 E(x_i^2 - 2x_i y_i + y_i^2) = \sum_{i=1}^d a_i^2 E(x_i - \frac{1}{2}x_i + \frac{1}{4}) \\ &= \sum_{i=1}^d a_i^2 (\frac{1}{2}E(x_i) + \frac{1}{4}) \end{aligned}$$

To maximize put all weight on the coordinate of x with highest probability of one. What if we used 1-norm instead of the two norm?

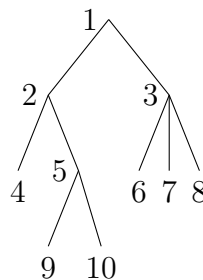
$$E(|x - y|) = E \sum_{i=1}^d a_i |x_i - y_i| = \sum_{i=1}^d a_i E |x_i - y_i| = \sum_{i=1}^d a_i b_i$$

where $b_i = E(x_i - y_i)$. If $\sum_{i=1}^d a_i^2 = 1$, then to maximize let $a_i = \frac{b_i}{b}$. Taking the dot product of a and b is maximized when both are in the same direction.

Solution: ■

Exercise 11.41 Maximize $x+y$ subject to the constraint that $x^2 + y^2 = 1$.

Exercise 11.42 Draw a tree with 10 vertices and label each vertex with a unique integer from 1 to 10. Construct the Prfer sequence for the tree. Given the Prfer sequence recreate the tree.



Solution:

The sequence is 2,3,3,1,5,5,2 ■

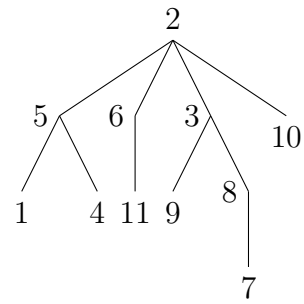
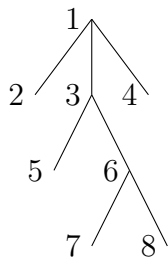
Exercise 11.43 Construct the tree corresponding to the following Prfer sequences

1. 113663

2. 552833226

Solution:

1. $(1,2), (1,4), (3,5), (6,7), (6,8), (1,3)$, and $(3,6)$
2. $(1,5), (4,5), (2,5), (7,8), (3,8), (3,9), (2,3), (2,10), (2,6)$, and $(6,11)$,



■

Index

- 2-norm, 119
- 2-universal, 237
- 4-way independence, 244
- Affinity matrix, 271
- Algorithm
 - greedy k -clustering, 262
 - k -means, 260
 - singular value decomposition, 117, 122
- Almost surely, 55
- Anchor term, 300
- Annulus, 12
- Aperiodic, 185
- Axioms
 - consistent, 286
 - for clustering, 286
 - rich, 286
 - scale invariant, 286
- Bad pair, 58
- Balanced k -means algorithm, 292
- Best fit, 6, 113
- Binomial distribution, 50
 - approximated by normal density, 50
 - approximated by Poisson, 52
- Boosting, 208
- Branching Process, 67
- Branching process, 70
- Breadth-first search, 64
- Cartesian coordinates, 13
- Characteristic equation, 371
- Characteristic function, 388
- Chebyshev's inequality, 10
- Clustering, 257
 - k -center criterion, 262
 - axioms, 286
 - balanced k -means algorithm, 292
 - k -means, 260
 - proper, 268
 - single link, 286
 - sparse cuts, 269
 - sum of pairs, 288
- CNF
 - CNF-sat, 83
- Cohesion, 277
- Commute time, 158
- Conductance, 151
- Coordinates
 - Cartesian, 13
 - polar, 13
- Coupon collector problem, 161
- Current
 - probabilistic interpretation, 154
- Cycles, 77
 - emergence, 76
 - number of, 76
- Data streams
 - counting frequent elements, 239
 - frequency moments, 234
 - frequent element, 240
 - majority element, 240
 - number of distinct elements, 235
 - number of occurrences of an element, 239
 - second moment, 241
- Degree distribution, 50
 - power law, 51
- Diagonalizable, 371
- Diameter of a graph, 57, 79
- Diameter two, 77
- Dimension reduction, 265
- Disappearance of isolated vertices, 77
- Discovery time, 156
- Distance
 - total variation, 171
- Distribution
 - vertex degree, 48
- Document ranking, 132
- Effective resistance, 158

- Eigenvalue, 371
- Eigenvector, 131, 371
- Electrical network, 151
- Equator
 - of sphere, 9, 16
- Erdős Rényi, 47
- Error correcting codes, 243
- Escape probability, 155
- Euler's constant, 162
- Expected degree
 - vertex, 47
- Exponential generating function, 387
- Extinct families
 - size, 74
- Extinction probability, 70, 72

- First moment method, 56
- Fourier transform, 337, 388
- Frequency domain, 338
- Frobenius norm, 119

- $G(n,p)$, 47
- Gamma function, 14
- Gaussian, 23, 388
 - annulus
 - width of, 24, 31
 - fitting to data, 26
- Gaussians
 - separating, 24
- Generating function, 71
 - component size, 91
 - for sum of two variables, 71
- Generating points on sphere, 22
- Giant component, 48, 55, 60, 62, 77
- Gibbs sampling, 173
- Graph
 - connectivity, 76
 - resistance, 161
- Graphical model, 307
- Greedy
 - k-clustering, 262
- Growth models, 87
 - nonuniform, 87
 - with preferential attachment, 96
 - without preferential attachment, 89

- Harmonic function, 151
- Hash function
 - universal, 237
- Heavy tail, 51
- Hidden Markov model, 302
- Hitting time, 156, 168

- Immortality probability, 72
- Incoherent, 338
- Increasing property, 55, 81
 - unsatisfiability, 83
- Independence
 - limited way, 243
- Indicator random variable, 58
 - of triangle, 53
- Intersection systems, 220
- Isolated vertices, 60, 77
 - number of, 60
- Isometry
 - restricted isometry property, 335

- Johnson-Lindenstrauss theorem, 33, 35

- k-center, 258
- k-clustering, 262
- k-means, 258
- k-means clustering algorithm, 260
- k-median, 258
- Kernel methods, 270
- Kleinberg, 98

- Law of large numbers, 9, 11
- Learning, 196
 - supervised, 271
 - unsupervised, 271
- Least squares, 113
- Linear separator, 198
- Linearity of expectation, 53
- Lloyd's algorithm, 260
- Local algorithm, 98
- Long-term probabilities, 149

- m-fold, 81

- Manifold
 - low dimensional, 271
- Margin, 198
 - maximum margin separator, 200
- Markov chain, 147
 - state, 170
- Markov Chain Monte Carlo, 148
- Markov random field, 310
- Markov's inequality, 10
- Matrix
 - multiplication
 - by sampling, 245
- Maximum cut problem, 133
- Maximum likelihood estimator, 27
- Maximum principle, 152
- MCMC, 148
- Metropolis-Hastings algorithm, 171
- Mixing time, 149
- Model
 - random graph, 47
- Molloy Reed, 88
- Moment generating function, 387
- Nearest neighbor, 272
- Nearest neighbor problem, 33, 36
- NMF, 299
- Nonnegative matrix factorization, 299
- Norm
 - 2-norm, 119
 - Frobenius, 119
- Normal distribution
 - standard deviation, 50
- Normalized conductance, 149, 177
- Number of triangles in $G(n, p)$, 53
- Orthonormal, 378
- Page rank, 166
 - personalized, 169
- Parallelepiped, 21
- Perceptron, 198
- Persistent, 147
- Phase transition, 55
 - CNF-sat, 83
 - nonfinite components, 94
- Polar coordinates, 13
- Polynomial interpolation, 243
- Power iteration, 132
- Power law distribution, 51
- Power method, 122
- Power-law distribution, 87
- Prüfer, 394
- Principle component analysis, 126
- Pseudo random, 244
- Pure-literal heuristic, 84
- Queue, 84
 - arrival rate, 85
- Radon, 216
- Random graph, 47
- Random projection, 33
 - theorem, 34
- Random walk
 - Eucleadean space, 162
 - in three dimensions, 163
 - in two dimensions, 163
 - on lattice, 162
 - undirected graph, 155
 - web, 166
- Rapid Mixing, 170
- Real spectral theorem, 373
- Recommendation system, 248
- Replication, 81
- Resistance, 151, 161
 - effective, 155
- Restart, 166
 - value, 166
- Return time, 166
- Sampling
 - length squared, 246
- Satisfying assignments
 - expected number of, 83
- Scale invariant, 286
- Second moment method, 53, 57
- Set system, 213, 217
- Sharp threshold, 55

- Shatter function, 217
- Shattered, 213
- Similar matrices, 371
- Similarity measure
 - cosine, 257
- Simplex, 22
- Single link, 286
- Singular value decomposition, 113
 - algorithm, 117
- Singular vector, 114
 - first, 114
 - left, 117
 - right, 117
 - second, 115
- Six-degrees separation, 98
- Sketch
 - matrix, 247
- Sketches
 - documents, 250
- Small world, 97
- Smallest-clause heuristic, 83
- Spam, 168
- Spectral clustering, 263
- Sphere
 - volume
 - narrow annulus, 19
 - near equator, 16
- Standard deviation
 - normal distribution, 50
- Stanley Milgram, 97
- State, 170
- Stationary distribution, 149
- Streaming model, 234
- Subgradient, 333
- Support vector, 202
- Support vector machine, 205
- Surface area
 - of sphere, 13
 - near equator, 20
- Symmetric matrices, 373
- Threshold, 55
 - CNF-sat, 83
- diameter $O(\ln n)$, 80
- disappearance of isolated vertices, 60
- emergence of cycles, 76
- emergence of diameter two, 57
- giant component plus isolated vertices, 78
- Time domain, 338
- Total variation distance, 171
- Trace, 381
- Triangles, 53
- Unit-clause heuristic, 84
- Unitary matrix, 378
- Unsatisfiability, 83
- Vapnik-Chervonenkis, 213
- VC theorem, 220
- VC-dimension, 210
 - convex polygons, 215
 - finite sets, 217
 - half spaces, 215
 - intervals, 214
 - pairs of intervals, 214
 - rectangles, 214
 - spheres, 216
- Vector space model, 6
- Vector space representation, 6
- Viterbi algorithm, 304
- Voltage
 - probabilistic interpretation, 153
- Volume
 - parallelepiped, 21
 - simplex, 22
 - sphere, 11
 - in narrow annulus, 19
 - near equator, 16
- Weak learner, 208
- World Wide Web, 166

References

- [ABC⁺08] Reid Andersen, Christian Borgs, Jennifer T. Chayes, John E. Hopcroft, Vahab S. Mirrokni, and Shang-Hua Teng. Local computation of pagerank contributions. *Internet Mathematics*, 5(1):23–45, 2008.
- [AF] David Aldous and James Fill. *Reversible Markov Chains and Random Walks on Graphs*. <http://www.stat.berkeley.edu/~aldous/RWG/book.html>.
- [AK] Sanjeev Arora and Ravindran Kannan. Learning mixtures of separated non-spherical gaussians. *Annals of Applied Probability*, 15(1A):6992.
- [Alo86] Noga Alon. Eigenvalues and expanders. *Combinatorica*, 6:83–96, 1986.
- [AM05] Dimitris Achlioptas and Frank McSherry. On spectral learning of mixtures of distributions. In *COLT*, pages 458–469, 2005.
- [AN72] Krishna Athreya and P. E. Ney. *Branching Processes*, volume 107. Springer, Berlin, 1972.
- [AP03] Dimitris Achlioptas and Yuval Peres. The threshold for random k-sat is $2^k (\ln 2 - o(k))$. In *STOC*, pages 223–231, 2003.
- [Aro11] Multiplicative weights method: a meta-algorithm and its applications. *Theory of Computing journal - to appear*, 2011.
- [AS08] Noga Alon and Joel H. Spencer. *The probabilistic method*. Wiley-Interscience Series in Discrete Mathematics and Optimization. John Wiley & Sons Inc., Hoboken, NJ, third edition, 2008. With an appendix on the life and work of Paul Erdős.
- [BA] Albert-Lszl Barabasi and Rka Albert. Emergence of scaling in random networks. *Science*, 286(5439).
- [BEHW] A. Blumer, A. Ehrenfeucht, D. Haussler, and M. K. Warmuth. Learnability and the vapnik-chervonenkis dimension. *Journal of the Association for Computing Machinery*.
- [Ble12] David M. Blei. Probabilistic topic models. *Commun. ACM*, 55(4):77–84, 2012.
- [BMPW98] Sergey Brin, Rajeev Motwani, Lawrence Page, and Terry Winograd. What can you do with a web in your pocket? *Data Engineering Bulletin*, 21:37–47, 1998.
- [Bol01] Béla Bollobás. *Random Graphs*. Cambridge University Press, 2001.

- [BT87] Béla Bollobás and Andrew Thomason. Threshold functions. *Combinatorica*, 7(1):35–38, 1987.
- [CF86] Ming-Te Chao and John V. Franco. Probabilistic analysis of two heuristics for the 3-satisfiability problem. *SIAM J. Comput.*, 15(4):1106–1118, 1986.
- [CGTS99] Moses Charikar, Sudipto Guha, Éva Tardos, and David B. Shmoys. A constant-factor approximation algorithm for the k-median problem (extended abstract). In *Proceedings of the thirty-first annual ACM symposium on Theory of computing*, STOC '99, pages 1–10, New York, NY, USA, 1999. ACM.
- [CHK⁺] Duncan S. Callaway, John E. Hopcroft, Jon M. Kleinberg, M. E. J. Newman, and Steven H. Strogatz. Are randomly grown graphs really random?
- [Chv92] *33rd Annual Symposium on Foundations of Computer Science, 24-27 October 1992, Pittsburgh, Pennsylvania, USA*. IEEE, 1992.
- [CLMW11] Emmanuel J. Candès, Xiaodong Li, Yi Ma, and John Wright. Robust principal component analysis? *J. ACM*, 58(3):11, 2011.
- [DFK91] Martin Dyer, Alan Frieze, and Ravindran Kannan. A random polynomial time algorithm for approximating the volume of convex bodies. *Journal of the Association for Computing Machinery*, 1991.
- [DFK⁺99] Petros Drineas, Alan M. Frieze, Ravi Kannan, Santosh Vempala, and V. Vinay. Clustering in large graphs and matrices. In *SODA*, pages 291–299, 1999.
- [DG99] Sanjoy Dasgupta and Anupam Gupta. An elementary proof of the johnson-lindenstrauss lemma. 99(006), 1999.
- [DS84] Peter G. Doyle and J. Laurie Snell. *Random walks and electric networks*, volume 22 of *Carus Mathematical Monographs*. Mathematical Association of America, Washington, DC, 1984.
- [DS07] Sanjoy Dasgupta and Leonard J. Schulman. A probabilistic analysis of em for mixtures of separated, spherical gaussians. *Journal of Machine Learning Research*, 8:203–226, 2007.
- [ER60] Paul Erdős and Alfred Rényi. On the evolution of random graphs. *Publication of the Mathematical Institute of the Hungarian Academy of Sciences*, 5:17–61, 1960.
- [Fel68] William Feller. *An Introduction to Probability Theory and Its Applications*, volume 1. Wiley, January 1968.
- [FK99] Alan M. Frieze and Ravindan Kannan. Quick approximation to matrices and applications. *Combinatorica*, 19(2):175–220, 1999.

- [Fri99] Friedgut. Sharp thresholds of graph properties and the k-sat problem. *Journal of the American Math. Soc.*, 12, no 4:1017–1054, 1999.
- [FS96] Alan M. Frieze and Stephen Suen. Analysis of two simple heuristics on a random instance of k-sat. *J. Algorithms*, 20(2):312–355, 1996.
- [GKP94] Ronald L. Graham, Donald E. Knuth, and Oren Patashnik. *Concrete mathematics - a foundation for computer science (2. ed.)*. Addison-Wesley, 1994.
- [GvL96] Gene H. Golub and Charles F. van Loan. *Matrix computations (3. ed.)*. Johns Hopkins University Press, 1996.
- [HBB10] Matthew D. Hoffman, David M. Blei, and Francis R. Bach. Online learning for latent dirichlet allocation. In *NIPS*, pages 856–864, 2010.
- [Jer98] Mark Jerrum. Mathematical foundations of the markov chain monte carlo method. In Dorit Hochbaum, editor, *Approximation Algorithms for NP-hard Problems*, 1998.
- [JKLP93] Svante Janson, Donald E. Knuth, Tomasz Łuczak, and Boris Pittel. The birth of the giant component. *Random Struct. Algorithms*, 4(3):233–359, 1993.
- [JLR00] Svante Janson, Tomasz Łuczak, and Andrzej Ruciński. *Random Graphs*. John Wiley and Sons, Inc, 2000.
- [Kan09] Ravindran Kannan. A new probability inequality using typical moments and concentration results. In *FOCS*, pages 211–220, 2009.
- [Kar90] Richard M. Karp. The transitive closure of a random digraph. *Random Structures and Algorithms*, 1(1):73–94, 1990.
- [Kle99] Jon M. Kleinberg. Authoritative sources in a hyperlinked environment. *JOURNAL OF THE ACM*, 46(5):604–632, 1999.
- [Kle00] Jon M. Kleinberg. The small-world phenomenon: an algorithm perspective. In *STOC*, pages 163–170, 2000.
- [Kle02] Jon M. Kleinberg. An impossibility theorem for clustering. In *NIPS*, pages 446–453, 2002.
- [KV95] Michael Kearns and Umesh Vazirani. *An introduction to Computational Learning Theory*. MIT Press, 1995.
- [KV09] Ravi Kannan and Santosh Vempala. Spectral algorithms. *Foundations and Trends in Theoretical Computer Science*, 4(3-4):157–288, 2009.
- [Liu01] Jun Liu. *Monte Carlo Strategies in Scientific Computing*. Springer, 2001.

- [Mat10] Jiří Matoušek. *Geometric discrepancy*, volume 18 of *Algorithms and Combinatorics*. Springer-Verlag, Berlin, 2010. An illustrated guide, Revised paperback reprint of the 1999 original.
- [Mit97] Tom M. Mitchell. *Machine Learning*. McGraw-Hill, New York, 1997.
- [MR95a] Michael Molloy and Bruce A. Reed. A critical point for random graphs with a given degree sequence. *Random Struct. Algorithms*, 6(2/3):161–180, 1995.
- [MR95b] Rajeev Motwani and Prabhakar Raghavan. *Randomized Algorithms*. Cambridge University Press, 1995.
- [MR99] Rajeev Motwani and Prabhakar Raghavan. Randomized algorithms. In *Algorithms and theory of computation handbook*, pages 15–1–15–23. CRC, Boca Raton, FL, 1999.
- [MU05] Michael Mitzenmacher and Eli Upfal. *Probability and computing - randomized algorithms and probabilistic analysis*. Cambridge University Press, 2005.
- [MV10] Ankur Moitra and Gregory Valiant. Settling the polynomial learnability of mixtures of gaussians. In *FOCS*, pages 93–102, 2010.
- [Pal85] Edgar M. Palmer. *Graphical evolution*. Wiley-Interscience Series in Discrete Mathematics. John Wiley & Sons Ltd., Chichester, 1985. An introduction to the theory of random graphs, A Wiley-Interscience Publication.
- [Par98] Beresford N. Parlett. *The symmetric eigenvalue problem*, volume 20 of *Classics in Applied Mathematics*. Society for Industrial and Applied Mathematics (SIAM), Philadelphia, PA, 1998. Corrected reprint of the 1980 original.
- [per10] *Markov Chains and Mixing Times*. American Mathematical Society, 2010.
- [Sch90] Rob Schapire. Strength of weak learnability. *Machine Learning*, 5:197–227, 1990.
- [SJ] Alistair Sinclair and Mark Jerrum. Approximate counting, uniform generation and rapidly mixing markov chains. *Information and Computation*.
- [Sly10] Allan Sly. Computational transition at the uniqueness threshold. In *FOCS*, pages 287–296, 2010.
- [SS01] Bernhard Scholkopf and Alexander J. Smola. *Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond*. MIT Press, Cambridge, MA, USA, 2001.
- [SWY75] G. Salton, A. Wong, and C. S. Yang. A vector space model for automatic indexing. *Commun. ACM*, 18:613–620, November 1975.

- [Val84] Leslie G. Valiant. A theory of the learnable. In *STOC*, pages 436–445, 1984.
- [Val13] L. Valiant. *Probably Approximately Correct: Nature’s Algorithms for Learning and Prospering in a Complex World*. Basic Books, 2013.
- [VC71] V. Vapnik and A. Chervonenkis. On the uniform convergence of relative frequencies of events to their probabilities. *Theory of Probability and its Applications*, 16(2):264–280, 1971.
- [Vem04] Santosh Vempala. *The Random Projection Method*. DIMACS, 2004.
- [VW02] Santosh Vempala and Grant Wang. A spectral algorithm for learning mixtures of distributions. *Journal of Computer and System Sciences*, pages 113–123, 2002.
- [Wil06] H.S. Wilf. *Generatingfunctionology*. Ak Peters Series. A K Peters, 2006.
- [WS98a] D. J. Watts and S. H. Strogatz. Collective dynamics of ‘small-world’ networks. *Nature*, 393 (6684), 1998.
- [WS98b] Duncan J. Watts and Steven H. Strogatz. Collective dynamics of ‘small-world’ networks. *Nature*, 393, 1998.
- [WW96] E. T. Whittaker and G. N. Watson. *A course of modern analysis*. Cambridge Mathematical Library. Cambridge University Press, Cambridge, 1996. An introduction to the general theory of infinite processes and of analytic functions; with an account of the principal transcendental functions, Reprint of the fourth (1927) edition.