

CS 6210: Matrix Computations

Introduction

David Bindel

2025-08-25

Logistics

CS 6210 is a graduate-level introduction to numerical linear algebra. We will study direct and iterative methods for linear systems, least squares problems, eigenvalue problems, and singular value problems. There is a detailed syllabus and rough schedule at the class web:

<http://www.cs.cornell.edu/courses/cs6210/2025fa/>

The web page is also your source for the homework, lecture notes, and course announcements. The web page also points to a discussion forum (Ed Discussions) and the Course Management System (CMS) used for homework submissions. There is also a link to a GitHub repository with the source for these notes. If you find an error, feel free to help me fix it!

I will assume that you have already had a course in linear algebra and that you know how to program. A previous numerical methods course is not required. The course will involve programming in Julia, so it will also be helpful – but not strictly necessary – for you to have prior exposure to MATLAB or Julia. If you know neither language but you have prior programming experience, you can learn them quickly enough.

A good general reference to follow throughout the course is the draft *Numerical Methods for Data Science* textbook. In particular, chapters 2 and 4 contain what I consider to be relevant material about how to program in Julia and about things you should know about linear algebra.

Linear algebra references

We assume some mathematical pre-requisites: linear algebra and “sufficient mathematical maturity.” This should be enough for most students, but some will want to brush up on their linear algebra. For those who want to review, or who want a reference while taking the course, I have a few recommendations.

There are a few options when it comes to linear algebra basics. At Cornell, our undergraduate linear algebra course uses the text by Lay (Lay, Lay, and McDonald 2016); the texts by Strang (Strang 2006, 2009) are a nice alternative. Strang's *Introduction to Linear Algebra* (Strang 2009) is the textbook for the MIT linear algebra course that is the basis for his enormously popular video lectures, available on MIT's OpenCourseWare site; if you prefer lecture to reading, Strang is known as an excellent lecturer.

These notes reflect the way that I teach matrix computations, but this is far from the only approach. The style of this class is probably closest to the style of the text by Demmel (Demmel 1997), which I recommend as a course text. The text of Trefethen and Bau (Trefethen and Bau 1997) covers roughly the same material, but with a different emphasis; I recommend it for an alternative perspective. At Cornell, our subscription to the SIAM book series means that both of these books are available as e-books to students on the campus network.

For students who intend to use matrix computations as a serious part of their future research career, the canonical reference is *Matrix Computations* by Golub and Van Loan (Golub and Van Loan 2013).

Basic notational conventions

In this section, we set out some basic notational conventions used in the class.

1. The complex unit is i (not i or j).
2. By default, all spaces in this class are finite dimensional. If there is only one space and the dimension is not otherwise stated, we use n to denote the dimension.
3. Concrete real and complex vector spaces are $\mathbb{R}^n$ and $\mathbb{C}^n$, respectively.
4. Real and complex matrix spaces are $\mathbb{R}^{m \times n}$ and $\mathbb{C}^{m \times n}$.
5. Unless otherwise stated, a concrete vector is a column vector.
6. The vector e is the vector of all ones.
7. The vector e_i has all zeros except a one in the i th place.
8. The basis $\{e_i\}_{i=1}^n$ is the *standard basis* in $\mathbb{R}^n$ or $\mathbb{C}^n$.
9. We use calligraphic math caps for abstract space, e.g. $\mathcal{U}, \mathcal{V}, \mathcal{W}$.
10. When we say U is a basis for a space $\mathcal{U}$, we mean U is an isomorphism $\mathcal{U} \rightarrow \mathbb{R}^n$. By a slight abuse of notation, we say U is a matrix whose columns are the abstract vectors $u_1, \dots, u_n$, and we write the linear combination $\sum_{i=1}^n u_i c_i$ concisely as Uc .
11. Similarly, $U^{-1}x$ represents the linear mapping from the abstract vector x to a concrete coefficient vector c such that $x = Uc$.

12. The space of univariate polynomials of degree at most d is $\mathcal{P}_d$.
13. Scalars will typically be lower case Greek, e.g. α, β . In some cases, we will also use lower case Roman letters, e.g. c, d .
14. Vectors (concrete or abstract) are denoted by lower case Roman, e.g. x, y, z .
15. Matrices and linear maps are both denoted by upper case Roman, e.g. A, B, C .
16. For $A \in \mathbb{R}^{m \times n}$, we denote the entry in row i and column j by a_{ij} . We reserve the notation A_{ij} to refer to a submatrix at block row i and block column j in a partitioning of A .
17. We use a superscript star to denote dual spaces and dual vectors; that is, $v^* \in \mathcal{V}^*$ is a dual vector in the space dual to $\mathcal{V}$.
18. In $\mathbb{R}^n$, we use x^* and x^T interchangeably for the transpose.
19. In $\mathbb{C}^n$, we use x^* and x^H interchangeably for the conjugate transpose.
20. Inner products are denoted by angles, e.g. $\langle x, y \rangle$. To denote an alternate inner product, we use subscripts, e.g. $\langle x, y \rangle_M = y^* M x$.
21. The standard inner product in $\mathbb{R}^n$ or $\mathbb{C}^n$ is also $x \cdot y$.
22. In abstract vector spaces with a standard inner product, we use v^* to denote the dual vector associated with v through the inner product, i.e. $v^* = (w \mapsto \langle w, v \rangle)$.
23. We use the notation $\|x\|$ to denote a norm of the vector x . As with inner products, we use subscripts to distinguish between multiple norms. When dealing with two generic norms, we will sometimes use the notation $\| \|y\| \|$ to distinguish the second norm from the first.
24. We use order notation for both algorithm scaling with parameters going to infinity (e.g. $O(n^3)$ time) and for reasoning about scaling with parameters going to zero (e.g. $O(\epsilon^2)$ error). We will rely on context to distinguish between the two.
25. We use *variational notation* to denote derivatives of matrix expressions, e.g. $\delta(AB) = \delta A B + A \delta B$ where δA and δB represent infinitesimal changes to the matrices A and B .
26. Symbols typeset in Courier font should be interpreted as code or pseudocode, e.g. $y = A * x$.
27. The function notation $\text{fl}(x)$ refers to taking a real or complex quantity (scalar or vector) and representing each entry in floating point.

Vectors

A vector space (or linear space) is a set of vectors that can be added or scaled in a sensible way – that is, addition is associative and commutative and scaling is distributive. We will generally denote vector spaces by script letters (e.g. $\mathcal{V}, \mathcal{W}$), vectors by lower case Roman letters (e.g. v, w), and scalars by lower case Greek letters (e.g. α, β). But we feel free to violate these conventions according to the dictates of our conscience or in deference to other conflicting conventions.

There are many types of vector spaces. Apart from the ubiquitous spaces $\mathbb{R}^n$ and $\mathbb{C}^n$, the most common vector spaces in applied mathematics are different types of function spaces. These include

$$\begin{aligned}\mathcal{P}_d &= \{\text{polynomials of degree at most } d\}; \\ \mathcal{V}^* &= \{\text{linear functions } \mathcal{V} \rightarrow \mathbb{R} \text{ (or } \mathbb{C})\}; \\ L(\mathcal{V}, \mathcal{W}) &= \{\text{linear maps } \mathcal{V} \rightarrow \mathcal{W}\}; \\ \mathcal{C}^k(\Omega) &= \{k\text{-times differentiable functions on a set } \Omega\};\end{aligned}$$

and many more. We compute with vectors in $\mathbb{R}^n$ and $\mathbb{C}^n$, which we represent concretely by tuples of numbers in memory, usually stored in sequence. To keep a broader perspective, though, we will also frequently describe examples involving the polynomial spaces $\mathcal{P}_d$.

Spanning sets and bases

We often think of a matrix as a set of vectors

$$V = \begin{bmatrix} v_1 & v_2 & \dots & v_n \end{bmatrix}.$$

The *range space* $\mathcal{R}(V)$ or the *span* $\text{sp}\{v_j\}_{j=1}^n$ is the set of vectors $\{Vc : c \in \mathbb{R}^n\}$ (or $\mathbb{C}^n$ for complex spaces). The vectors are *linearly independent* if any vector in the span has a unique representation as a linear combination of the spanning vectors; equivalently, the vectors are linearly independent if there is no linear combination of the spanning vectors that equals zero except the one in which all coefficients are zero. A linearly independent set of vectors is a *basis* for a space $\mathcal{V}$ if $\text{sp}\{v_j\}_{j=1}^n = \mathcal{V}$; the number of basis vectors n is the *dimension* of the space. Spaces like $\mathcal{C}^k(\Omega)$ do not generally have a finite basis; they are *infinite-dimensional*. We will focus on finite-dimensional spaces in the class, but it is useful to know that there are interesting infinite-dimensional spaces in the broader world.

The *standard basis* in $\mathbb{R}^n$ is the set of column vectors e_j with a one in the j th position and zeros elsewhere:

$$I = \begin{bmatrix} e_1 & e_2 & \dots & e_n \end{bmatrix}.$$

Of course, this is not the only basis. Any other set of n linearly independent vectors in $\mathbb{R}^n$ is also a basis

$$V = \begin{bmatrix} v_1 & v_2 & \dots & v_n \end{bmatrix}.$$

The matrix V formed in this way is invertible, with multiplication by V corresponding to a change from the $\{v_j\}$ basis to the standard basis and multiplication by V^{-1} corresponding to a map from the standard basis into the $\{v_j\}$ basis.

We will use matrix-like notation to describe bases and spanning sets even when we deal with spaces other than $\mathbb{R}^n$. For example, a common basis for $\mathcal{P}_d$ is the *power basis*

$$X = \begin{bmatrix} 1 & x & x^2 & \dots & x^d \end{bmatrix}.$$

Each “column” in X is really a function of one variable (x), and matrix-vector multiplication with X represents a map from a coefficient vector in $\mathbb{R}^{d+1}$ to a polynomial in $\mathcal{P}_d$. That is, we write polynomials $p \in \mathcal{P}_d$ in terms of the basis as

$$p = Xc = \sum_{j=0}^d c_j x^j$$

and, we think of computing the coefficients from the abstract polynomial via a formal inverse:

$$c = X^{-1}p.$$

We typically think of a map like $Y^* = X^{-1}$ in terms of “rows”

$$Y^* = \begin{bmatrix} y_0^* \\ y_1^* \\ \vdots \\ y_d^* \end{bmatrix}$$

where each row y_j^* is a *linear functional* or *dual vector* (i.e. linear mappings from the vector space to the real or complex numbers). Collectively, $\{y_0^*, \dots, y_d^*\}$ are the *dual basis* to $\{1, x, \dots, x^d\}$.

The power basis is not the only basis for $\mathcal{P}_d$. Other common choices include Newton or Lagrange polynomials with respect to a set of points, which you may have seen in another class such as CS 4210. In this class, we will sometimes use the Chebyshev¹ polynomial basis $\{T_j(x)\}$ given by the recurrence

$$\begin{aligned} T_0(x) &= 1 \\ T_1(x) &= x \\ T_{j+1}(x) &= 2xT_j(x) - T_{j-1}(x), \quad j \geq 1, \end{aligned}$$

and the Legendre polynomial basis $\{P_j(x)\}$, given by the recurrence

$$\begin{aligned} P_0(x) &= 1 \\ P_1(x) &= x \\ (j+1)P_{j+1}(x) &= (2j+1)xP_j(x) - jP_{j-1}(x). \end{aligned}$$

¹Pafnuty Chebyshev was a nineteenth century Russian mathematician, and his name has been transliterated from the Cyrillic alphabet into the Latin alphabet in several different ways. We inherit our usual spelling from one of the French transliterations, but the symbol T for the polynomials comes from the German transliteration Tschebyscheff.

As we will see over the course of the semester, sometimes the “obvious” choice of basis (e.g. the standard basis in $\mathbb{R}^n$ or the power basis in $\mathcal{P}_d$) is not the best choice for numerical computations.

Vector norms

A *norm* $\|\cdot\|$ measures vector lengths. It is positive definite, homogeneous, and sub-additive:

$$\begin{aligned}\|v\| &\geq 0 \text{ and } \|v\| = 0 \text{ iff } v = 0 \\ \|\alpha v\| &= |\alpha| \|v\| \\ \|u + v\| &\leq \|u\| + \|v\|.\end{aligned}$$

The three most common vector norms we work with in $\mathbb{R}^n$ are the Euclidean norm (aka the 2-norm), the ∞ -norm (or max norm), and the 1-norm:

$$\begin{aligned}\|v\|_2 &= \sqrt{\sum_j |v_j|^2} \\ \|v\|_\infty &= \max_j |v_j| \\ \|v\|_1 &= \sum_j |v_j|\end{aligned}$$

Many other norms can be related to one of these three norms. In particular, a “natural” norm in an abstract vector space will often look strange in the corresponding concrete representation with respect to some basis function. For example, consider the vector space of polynomials with degree at most 2 on $[-1, 1]$. This space also has a natural Euclidean norm, max norm, and 1-norm; for a given polynomial $p(x)$ these are

$$\begin{aligned}\|p\|_2 &= \sqrt{\int_{-1}^1 |p(x)|^2 dx} \\ \|p\|_\infty &= \max_{x \in [-1, 1]} |p(x)| \\ \|p\|_1 &= \int_{-1}^1 |p(x)| dx.\end{aligned}$$

But when we write $p(x)$ in terms of the coefficient vector with respect to the power basis (for example), the max norm of the polynomial is not the same as the max norm of the coefficient vector. In fact, if we consider a polynomial $p(x) = c_0 + c_1x$, then the max norm of the polynomial p is the same as the one-norm of the coefficient vector — the proof of which is left as a useful exercise to the reader.

In a finite-dimensional vector space, all norms are *equivalent*: that is, if $\|\cdot\|$ and $|||\cdot|||$ are two norms on the same finite-dimensional vector space, then there exist constants c and C such that for any v in the space,

$$c\|v\| \leq |||v||| \leq C\|v\|.$$

Of course, there is no guarantee that the constants are small!

An *isometry* is a mapping that preserves vector norms. For $\mathbb{R}^n$, the only isometries for the 1-norm and the ∞ -norm are permutations. For Euclidean space, though, there is a much richer set of isometries, represented by the orthogonal matrices (matrices s.t. $Q^*Q = I$).

Inner products

An *inner product* $\langle \cdot, \cdot \rangle$ is a function from two vectors into the real numbers (or complex numbers for a complex vector space). It is positive definite, linear in the first slot, and symmetric (or Hermitian in the case of complex vectors); that is:

$$\begin{aligned}\langle v, v \rangle &\geq 0 \text{ and } \langle v, v \rangle = 0 \text{ iff } v = 0 \\ \langle \alpha u, w \rangle &= \alpha \langle u, w \rangle \text{ and } \langle u + v, w \rangle = \langle u, w \rangle + \langle v, w \rangle \\ \langle u, v \rangle &= \overline{\langle v, u \rangle},\end{aligned}$$

where the overbar in the latter case corresponds to complex conjugation. A vector space with an inner product is sometimes called an *inner product space* or a *Euclidean space*.

Every inner product defines a corresponding norm

$$\|v\| = \sqrt{\langle v, v \rangle}$$

The inner product and the associated norm satisfy the *Cauchy-Schwarz* inequality

$$\langle u, v \rangle \leq \|u\| \|v\|.$$

The *standard inner product* on $\mathbb{C}^n$ is

$$x \cdot y = y^* x = \sum_{j=1}^n \bar{y}_j x_j.$$

But the standard inner product is not the only inner product, just as the standard Euclidean norm is not the only norm.

Just as norms allow us to reason about size, inner products let us reason about angles. In particular, we define the cosine of the angle θ between nonzero vectors v and w as

$$\cos \theta = \frac{\langle v, w \rangle}{\|v\| \|w\|}.$$

Returning to our example of a vector space of polynomials, the standard $L^2([-1, 1])$ inner product is

$$\langle p, q \rangle_{L^2([-1, 1])} = \int_{-1}^1 p(x) \bar{q}(x) dx.$$

If we express p and q with respect to a basis (e.g. the power basis), we find that we can represent this inner product via a symmetric positive definite matrix. For example, let $p(x) = c_0 + c_1x + c_2x^2$ and let $q(x) = d_0 + d_1x + d_2x^2$. Then

$$\langle p, q \rangle_{L^2([-1,1])} = \begin{bmatrix} d_0 \\ d_1 \\ d_2 \end{bmatrix}^* \begin{bmatrix} a_{00} & a_{01} & a_{02} \\ a_{10} & a_{11} & a_{12} \\ a_{20} & a_{21} & a_{22} \end{bmatrix} \begin{bmatrix} c_0 \\ c_1 \\ c_2 \end{bmatrix} = d^* A c = \langle c, d \rangle_A$$

where

$$a_{ij} = \int_{-1}^1 x^{i-1} x^{j-1} dx = \begin{cases} 2/(i+j-1), & i+j \text{ even} \\ 0, & \text{otherwise} \end{cases}$$

The symmetric positive definite matrix A is what is sometimes called the *Gram matrix* for the basis $\{1, x, x^2\}$.

We say two vectors u and v are *orthogonal* with respect to an inner product if $\langle u, v \rangle = 0$. If u and v are orthogonal, we have the *Pythagorean theorem*:

$$\|u + v\|^2 = \langle u + v, u + v \rangle = \langle u, u \rangle + \langle v, u \rangle + \langle u, v \rangle + \langle v, v \rangle = \|u\|^2 + \|v\|^2$$

Two vectors u and v are *orthonormal* if they are orthogonal with respect to the inner product and have unit length in the associated norm. When we work in an inner product space, we often use an *orthonormal basis*, i.e. a basis in which all the vectors are orthonormal. For example, the normalized Legendre polynomials

$$\sqrt{\frac{2}{2j+1}} P_j(x)$$

form orthonormal bases for the $\mathcal{P}_d$ inner product with respect to the $L^2([-1, 1])$ inner product.

Matrices and mappings

A matrix represents a mapping between two vector spaces. That is, if $L : \mathcal{V} \rightarrow \mathcal{W}$ is a linear map, then the associated matrix A with respect to bases V and W satisfies $A = W^{-1}LV$. The same linear mapping corresponds to different matrices depending on the choices of basis. But matrices can represent several other types of mappings as well. Over the course of this class, we will see several interpretations of matrices:

- **Linear maps.** A map $L : \mathcal{V} \rightarrow \mathcal{W}$ is linear if $L(x + y) = Lx + Ly$ and $L(\alpha x) = \alpha Lx$. The corresponding matrix is $A = W^{-1}LV$.
- **Linear operators.** A linear map from a space to itself ($L : \mathcal{V} \rightarrow \mathcal{V}$) is a linear operator. The corresponding (square) matrix is $A = V^{-1}LV$.

- **Bilinear forms.** A map $a : \mathcal{V} \times \mathcal{W} \rightarrow \mathbb{R}$ (or $\mathbb{C}$ for complex spaces) is bilinear if it is linear in both slots: $a(\alpha u + v, w) = \alpha a(u, w) + a(v, w)$ and $a(v, \alpha u + w) = \alpha a(v, u) + a(v, w)$. The corresponding matrix has elements $A_{ij} = a(v_i, w_j)$; if $v = Vc$ and $w = Wd$ then $a(v, w) = d^T A c$.

We call a bilinear form on $\mathcal{V} \times \mathcal{V}$ *symmetric* if $a(v, w) = a(w, v)$; in this case, the corresponding matrix A is also symmetric ($A = A^T$). A symmetric form and the corresponding matrix are called *positive semi-definite* if $a(v, v) \geq 0$ for all v . The form and matrix are *positive definite* if $a(v, v) > 0$ for any $v \neq 0$.

A *skew-symmetric* matrix ($A = -A^T$) corresponds to a skew-symmetric or anti-symmetric bilinear form, i.e. $a(v, w) = -a(w, v)$.

- **Sesquilinear forms.** A map $a : \mathcal{V} \times \mathcal{W} \rightarrow \mathbb{C}$ (where $\mathcal{V}$ and $\mathcal{W}$ are complex vector spaces) is sesquilinear if it is linear in the first slot and the conjugate is linear in the second slot: $a(\alpha u + v, w) = \alpha a(u, w) + a(v, w)$ and $a(v, \alpha u + w) = \bar{\alpha} a(v, u) + a(v, w)$. The matrix has elements $A_{ij} = a(v_i, w_j)$; if $v = Vc$ and $w = Wd$ then $a(v, w) = d^* A c$.

We call a sesquilinear form on $\mathcal{V} \times \mathcal{V}$ *Hermitian* if $a(v, w) = \overline{a(w, v)}$; in this case, the corresponding matrix A is also Hermitian ($A = A^*$). A Hermitian form and the corresponding matrix are called *positive semi-definite* if $a(v, v) \geq 0$ for all v . The form and matrix are *positive definite* if $a(v, v) > 0$ for any $v \neq 0$.

A *skew-Hermitian* matrix ($A = -A^*$) corresponds to a skew-Hermitian or anti-Hermitian bilinear form, i.e. $a(v, w) = -\overline{a(w, v)}$.

- **Quadratic forms.** A quadratic form $\phi : \mathcal{V} \rightarrow \mathbb{R}$ (or $\mathbb{C}$) is a homogeneous quadratic function on $\mathcal{V}$, i.e. $\phi(\alpha v) = \alpha^2 \phi(v)$ for which the map $b(v, w) = \phi(v + w) - \phi(v) - \phi(w)$ is bilinear. Any quadratic form on a finite-dimensional space can be represented as $c^* A c$ where c is the coefficient vector for some Hermitian matrix A . The formula for the elements of A given ϕ is left as an exercise.

We care about linear maps and linear operators almost everywhere, and most students come out of a first linear algebra class with some notion that these are important. But apart from very standard examples (inner products and norms), many students have only a vague notion of what a bilinear form, sesquilinear form, or quadratic form might be. Bilinear forms and sesquilinear forms show up when we discuss large-scale solvers based on projection methods. Quadratic forms are important in optimization, physics (where they often represent energy), and statistics (e.g. for understanding variance and covariance).

Matrix norms

The space of matrices forms a vector space; and, as with other vector spaces, it makes sense to talk about norms. In particular, we frequently want norms that are *consistent* with vector

norms on the range and domain spaces; that is, for any w and v , we want

$$w = Av \implies \|w\| \leq \|A\|\|v\|.$$

One “obvious” consistent norm is the *Frobenius norm*,

$$\|A\|_F = \sqrt{\sum_{i,j} a_{ij}^2}.$$

Even more useful are *induced norms* (or *operator norms*)

$$\|A\| = \sup_{v \neq 0} \frac{\|Av\|}{\|v\|} = \sup_{\|v\|=1} \|Av\|.$$

The induced norms corresponding to the vector 1-norm and ∞ -norm are

$$\begin{aligned} \|A\|_1 &= \max_j \sum_i |a_{ij}| \quad (\text{max column sum}) \\ \|A\|_\infty &= \max_i \sum_j |a_{ij}| \quad (\text{max row sum}) \end{aligned}$$

The norm induced by the vector Euclidean norm (variously called the matrix 2-norm or the spectral norm) is more complicated.

The Frobenius norm and the matrix 2-norm are both *orthogonally invariant* (or *unitarily invariant* in a complex vector space). That is, if Q is a square matrix with $Q^* = Q^{-1}$ (an orthogonal or unitary matrix) of the appropriate dimensions

$$\begin{aligned} \|QA\|_F &= \|A\|_F, & \|AQ\|_F &= \|A\|_F, \\ \|QA\|_2 &= \|A\|_2, & \|AQ\|_2 &= \|A\|_2. \end{aligned}$$

This property will turn out to be frequently useful throughout the course.

Decompositions and canonical forms

Matrix decompositions (also known as *matrix factorizations*) are central to numerical linear algebra. We will get to know six such factorizations well:

- $PA = LU$ (a.k.a. Gaussian elimination). Here L is unit lower triangular (triangular with 1 along the main diagonal), U is upper triangular, and P is a permutation matrix.
- $A = LL^*$ (a.k.a. Cholesky factorization). Here A is Hermitian and positive definite, and L is a lower triangular matrix.
- $A = QR$ (a.k.a. QR decomposition). Here Q has orthonormal columns and R is upper triangular. If we think of the columns of A as a basis, QR decomposition corresponds to the Gram-Schmidt orthogonalization process you have likely seen in the past (though we rarely compute with Gram-Schmidt).

- $A = U\Sigma V^*$ (a.k.a. the singular value decomposition or SVD). Here U and V have orthonormal columns and Σ is diagonal with non-negative entries.
- $A = Q\Lambda Q^*$ (a.k.a. symmetric eigendecomposition). Here A is Hermitian (symmetric in the real case), Q is orthogonal or unitary, and Λ is a diagonal matrix with real numbers on the diagonal.
- $A = QTQ^*$ (a.k.a. Schur form). Here A is a square matrix, Q is orthogonal or unitary, and T is upper triangular (or nearly so).

The last three of these decompositions correspond to *canonical forms* for abstract operators. That is, we can view these decompositions as finding bases in which the matrix representation of some operator or form is particularly simple. In a first linear algebra course, one generally considers canonical forms associated with general bases (not restricted to be orthogonal):

- For a linear map, we have the canonical form

$$A = U^{-1}AV = \begin{bmatrix} I_k & 0 \\ 0 & 0 \end{bmatrix}$$

where k is the rank of the matrix and the zero blocks are sized so the dimensions make sense.

- For an operator, we have the *Jordan canonical form*,

$$J = V^{-1}AV = \begin{bmatrix} J_{\lambda_1} & & \\ & J_{\lambda_2} & \\ & & \ddots & J_{\lambda_r} \end{bmatrix}$$

where each J_λ is a Jordan block with λ down the main diagonal and 1 on the first superdiagonal.

- For a quadratic form, we have the canonical form

$$a(Vx) = \sum_{i=1}^{k_+} x_i^2 - \sum_{i=k_++1}^{k_++k_-} x_i^2 = x^T Ax, \quad A = \begin{bmatrix} I_{k_+} & & \\ & -I_{k_-} & \\ & & 0_{k_0} \end{bmatrix}.$$

The integer triple (k_+, k_0, k_-) is sometimes called the *inertia* of the quadratic form (or *Sylvester's inertia*).

As beautiful as these canonical forms are, they are terrible for computation. In general, they need not even be continuous! However, if V and U have inner products, it makes sense to restrict our attention to orthonormal bases. This restriction gives canonical forms that we tend to prefer in practice:

- For a linear map, we have the canonical form

$$U^{-1}AV = \begin{bmatrix} \Sigma_k & 0 \\ 0 & 0 \end{bmatrix}$$

where k is the rank of the matrix and the zero blocks are sized so the dimensions make sense. The matrix Σ_k is a diagonal matrix of *singular values*

$$\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_k > 0,$$

and the bases U and V consist of the *singular vectors*.

- For an operator, we have the *Schur canonical form*,

$$V^{-1}AV = T$$

where T is an upper triangular matrix (in the complex case) or a quasi-upper triangular matrix that may have 2-by-2 blocks (in the case of a real matrix with complex eigenvalues). In this case, the basis vectors span nested invariant subspaces of $\mathcal{A}$.

- For a quadratic form, we have the canonical form

$$a(Vx) = \sum_{i=1}^n \lambda_i x_i^2 = x^T \Lambda x,$$

where Λ is a diagonal matrix with $\lambda_1, \dots, \lambda_n$ on the diagonal.

The SVD and the 2-norm

The singular value decomposition is useful for a variety of reasons; we close off the lecture by showing one such use.

Suppose $A = U\Sigma V^*$ is the singular value decomposition of some matrix. Using orthogonal invariance (unitary invariance) of the 2-norm, we have

$$\|A\|_2 = \|U^*AV\|_2 = \|\Sigma\|_2,$$

i.e.

$$\|A\|_2 = \max_{\|v\|_2=1} \frac{\sum_j \sigma_j |v_j|^2}{\sum |v_j|^2}.$$

That is, the spectral norm is the largest weighted average of the singular values, which is the same as just the largest singular value.

The small singular values also have a meaning. If A is a square, invertible matrix then

$$\|A^{-1}\|_2 = \|V\Sigma^{-1}U^*\|_2 = \|\Sigma^{-1}\|_2,$$

i.e. $\|A^{-1}\|_2$ is the inverse of the smallest singular value of A .

The smallest singular value of a nonsingular matrix A can also be interpreted as the “distance to singularity”: if σ_n is the smallest singular value of A , then there is a matrix E such that $\|E\|_2 = \sigma_n$ and $A + E$ is singular; and there is no such matrix with smaller norm.

These facts about the singular value decomposition are worth pondering, as they will be particularly useful in the next lecture when we ponder sensitivity and conditioning.

References

- Demmel, James. 1997. *Applied Numerical Linear Algebra*. SIAM.
- Golub, Gene, and Charles Van Loan. 2013. *Matrix Computations*. Fourth. Johns Hopkins University Press.
- Lay, David, Steven Lay, and Judi McDonald. 2016. *Linear Algebra and Its Applications*. Fifth. Pearson.
- Strang, Gilbert. 2006. *Linear Algebra and Its Applications*. Fourth. Brooks/Cole Publishing.
- . 2009. *Introduction to Linear Algebra*. Fourth. Wellesley-Cambridge Press.
- Trefethen, Lloyd N, and David Bau III. 1997. *Numerical Linear Algebra*. SIAM.