

CS 6210: Matrix Computations

Power method and subspace iteration

David Bindel

2025-10-27

Power iteration

In most introductory linear algebra classes, one computes eigenvalues as roots of a characteristic polynomial. For most problems, this is a *bad idea*: the roots of the characteristic polynomial are often very sensitive to changes in the polynomial coefficients even when they correspond to well-conditioned eigenvalues. Rather than starting from this point, we will start with another idea: the *power iteration*.

Suppose $A \in \mathbb{C}^{n \times n}$ is diagonalizable, with eigenvalues $\lambda_1, \dots, \lambda_n$ ordered so that

$$|\lambda_1| \geq |\lambda_2| \geq \dots \geq |\lambda_n|.$$

Then we have $A = V\Lambda V^{-1}$ where $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_n)$. Now, note that

$$A^k = (V\Lambda V^{-1})(V\Lambda V^{-1}) \dots (V\Lambda V^{-1}) = V\Lambda^k V^{-1},$$

or, to put it differently,

$$A^k V = V \Lambda^k.$$

Now, suppose we have a randomly chosen vector $x = V\tilde{x} \in \mathbb{C}^n$, and consider

$$A^k x = A^k V\tilde{x} = V\Lambda^k \tilde{x} = \sum_{j=1}^n v_j \lambda_j^k \tilde{x}_j.$$

If we pull out a constant factor from this expression, we have

$$A^k x = \lambda_1^k \left(\sum_{j=1}^n v_j \left(\frac{\lambda_j}{\lambda_1} \right)^k \tilde{x}_j \right).$$

If $|\lambda_1| > |\lambda_2|$, then $(\lambda_j/\lambda_1)^k \rightarrow 0$ for each $j > 1$, and for large enough k , we expect $A^k x$ to be nearly parallel to v_1 , assuming $\tilde{x}_1 \neq 0$. This is the idea behind the power iteration:

$$x^{(k+1)} = \frac{Ax^{(k)}}{\|Ax^{(k)}\|} = \frac{A^k x^{(0)}}{\|A^k x^{(0)}\|}.$$

Assuming that the first component of $V^{-1}x^{(0)}$ is nonzero and that $|\lambda_1| > |\lambda_2|$, the iterates $x^{(k)}$ converge linearly to the “dominant” eigenvector of A , with the error asymptotically decreasing by a factor of $|\lambda_1|/|\lambda_2|$ at each step.

There are three obvious potential problems with the power method:

1. What if the first component of $V^{-1}x^{(0)}$ is zero?
2. What if λ_1/λ_2 is near one?
3. What if we want the eigenpair (λ_j, v_j) for $j \neq 1$?

The first point turns out to be a non-issue: if we choose $x^{(0)}$ at random, then the first component of $V^{-1}x^{(0)}$ will be nonzero with probability 1. Even if we were so extraordinarily unlucky as to choose a starting vector for which $V^{-1}x^{(0)}$ *did* have a zero leading coefficient, perturbations due to floating point arithmetic would generally bump us to the case in which we had a nonzero coefficient.

The second and third points turn out to be more interesting, and we address them now.

Spectral transformation and shift-invert

Suppose again that A is diagonalizable with $A = V\Lambda V^{-1}$. The power iteration relies on the identity

$$A^k = V\Lambda^k V^{-1}.$$

Now, suppose that $f(z)$ is any function that is defined locally by a convergent power series. Then as long as the eigenvalues are within the radius of convergence, we can define $f(A)$ via the same power series, and

$$f(A) = Vf(\Lambda)V^{-1}$$

where $f(\Lambda) = \text{diag}(f(\lambda_1), f(\lambda_2), \dots, f(\lambda_n))$. So the spectrum of $f(A)$ is the image of the spectrum of A under the mapping f , a fact known as the *spectral mapping theorem*.

As a particular instance, consider the function $f(z) = (z - \sigma)^{-1}$. This gives us

$$(A - \sigma I)^{-1} = V(\Lambda - \sigma I)^{-1}V^{-1},$$

and so if we run power iteration on $(A - \sigma I)^{-1}$, we will converge to the eigenvector corresponding to the eigenvalue λ_j for which $(\lambda_j - \sigma)^{-1}$ is maximal — that is, we find the eigenvalue closest to σ in the complex plane. Running the power method on $(A - \sigma I)^{-1}$ is sometimes called the shift-invert power method.

Changing shifts

If we know a shift σ that is close to a desired eigenvalue, the shift-invert power method may be a reasonable method. But even with a good choice of shift, this method converges at best linearly (i.e. the error goes down by a constant factor at each step). We can do better by choosing a shift *dynamically*, so that as we improve the eigenvector, we also get a more accurate shift.

Suppose $\hat{v}$ is an approximate eigenvector for A , i.e. we can find some $\hat{\lambda}$ so that

$$A\hat{v} - \hat{v}\hat{\lambda} \approx 0. \quad (1)$$

The choice of corresponding approximate eigenvalues is not so clear, but a reasonable choice (which is always well-defined when $\hat{v}$ is nonzero) comes from multiplying Equation 1 by $\hat{v}^*$ and changing the $\approx$ to an equal sign:

$$\hat{v}^* A \hat{v} - \hat{v}^* \hat{v} \hat{\lambda} = 0.$$

The resulting eigenvalue approximation $\hat{\lambda}$ is the *Rayleigh quotient*:

$$\hat{\lambda} = \frac{\hat{v}^* A \hat{v}}{\hat{v}^* \hat{v}}.$$

If we dynamically choose shifts for shift-invert steps using Rayleigh quotients, we get the *Rayleigh quotient iteration*:

$$\begin{aligned} \lambda_{k+1} &= \frac{v^{(k)*} A v^{(k)}}{v^{(k)*} v^{(k)}} \\ v^{(k+1)} &= \frac{(A - \lambda_{k+1})^{-1} v^{(k)}}{\|(A - \lambda_{k+1})^{-1} v^{(k)}\|_2} \end{aligned}$$

Unlike the power method, the Rayleigh quotient iteration has locally quadratic convergence — so once convergence sets in, the number of correct digits roughly doubles from step to step. We will return to this method later when we discuss symmetric matrices, for which the Rayleigh quotient iteration has locally *cubic* convergence.

Subspaces and orthogonal iteration

So far, we have still not really addressed the issue of dealing with clustered eigenvalues. For example, in power iteration, what should we do if λ_1 and λ_2 are very close? If the ratio between the two eigenvalues is nearly one, we don't expect the power method to converge quickly; and we are likely to not have at hand a shift which is much closer to λ_1 than to λ_2 , so shift-invert power iteration might not help much. In this case, we might want to relax our question, and look for the invariant subspace associated with λ_1 and λ_2 (and maybe more eigenvalues if there are more of them clustered together with λ_1) rather than looking for the eigenvector associated with λ_1 . This is the idea behind *subspace iteration*.

In subspace iteration, rather than looking at $A^k x_0$ for some initial vector x_0 , we look at $\mathcal{V}_k = A^k \mathcal{V}_0$, where $\mathcal{V}_0$ is some initial subspace. If $\mathcal{V}_0$ is a p -dimensional space, then under some mild assumptions the space $\mathcal{V}_k$ will asymptotically converge to the p -dimensional invariant subspace of A associated with the p eigenvalues of A with largest modulus. The analysis is basically the same as the analysis for the power method. In order to actually *compute*, though, we need bases for the subspaces $\mathcal{V}_k$. Let us define these bases by the recurrence

$$Q_{k+1} R_{k+1} = A Q_k$$

where Q_0 is a matrix with p orthonormal columns and $Q_{k+1} R_{k+1}$ represents an economy QR decomposition. This recurrence is called *orthogonal iteration*, since the columns of Q_{k+1} are an orthonormal basis for the range space of $A Q_k$, and the span of Q_k is the span of $A^k Q_0$.

Assuming there is a gap between $|\lambda_p|$ and $|\lambda_{p+1}|$, orthogonal iteration will usually converge to an orthonormal basis for the invariant subspace spanned by the first p eigenvectors of A . But it is interesting to look not only at the behavior of the subspace, but also at the span of the individual eigenvectors. For example, notice that the first column $q_{k,1}$ of Q_k satisfies the recurrence

$$q_{k+1,1} r_{k+1,11} = A q_{k,1},$$

which means that the vectors $q_{k,1}$ evolve according to the power method! So over time, we expect the first columns of the Q_k to converge to the dominant eigenvector. Similarly, we expect the first two columns of Q_k to converge to a basis for the dominant two-dimensional invariant subspace, the first three columns to converge to the dominant three-dimensional invariant subspace, and so on. This observation suggests that we might be able to get a complete list of nested invariant subspaces by letting the initial Q_0 be some square matrix. This is the basis for the workhorse of nonsymmetric eigenvalue algorithms, the *QR method*, which we will turn to next time.