

Iterative solvers for linear equations

Daniel A. Spielman

October 21, 2009

15.1 Overview

In this and the next lecture, I will discuss iterative algorithms for solving linear equations in Laplacian matrices and symmetric diagonally dominant matrices. I will begin with an introduction to iterative linear solvers in general, and then finish by explaining how this is relevant to Laplacians.

15.2 Why iterative methods?

One is first taught to solve linear systems like

$$Ax = b$$

by *direct methods* such as Gaussian elimination, computing the inverse of A , or the LU factorization. However, all of these algorithms can be very slow. This is especially true when A is sparse. Just writing down the inverse takes $O(n^2)$ space, and computing the inverse takes $O(n^3)$ time if we do it naively. This might be OK if A is dense. But, it is very wasteful if A only has $O(n)$ non-zero entries.

In general, we prefer algorithms whose running time is proportional to the number of non-zero entries in the matrix A , and which do not require much more space than that used to store A .

Iterative algorithms solve linear equations *while only performing multiplications* by A , and performing a few vector operations. Unlike the *direct methods* which are based on elimination, the iterative algorithms do not get exact solutions. Rather, they get closer and closer to the solution the longer they work. The advantage of these methods is that they need to store very little, and are often much faster than the direct methods. When A is symmetric, the running times of these methods are determined by the eigenvalues of A .

15.3 First-Order Richardson Iteration

To get started, we will examine a simple, but sub-optimal, iterative method, Richardson's iteration. The idea of the method is to find an iterative process that has the solution to $Ax = b$ as a fixed

point, and which converges. We observe that if $A\mathbf{x} = \mathbf{b}$, then for any α ,

$$\begin{aligned}\alpha A\mathbf{x} &= \alpha \mathbf{b}, & \implies \\ \mathbf{x} + (\alpha A - I)\mathbf{x} &= \alpha \mathbf{b}, & \implies \\ \mathbf{x} &= (I - \alpha A)\mathbf{x} + \alpha \mathbf{b}.\end{aligned}$$

This leads us to the following iterative process:

$$\mathbf{x}^t = (I - \alpha A)\mathbf{x}^{t-1} + \alpha \mathbf{b},$$

where we will take $\mathbf{x}^0 = \mathbf{0}$. We will show that this converges if

$$I - \alpha A$$

has norm less than 1, and that the convergence rate depends on how much the norm is less than 1. As we are assuming A is symmetric, $I - \alpha A$ is symmetric as well, and so its norm is the maximum absolute value of its eigenvalues. Let $\lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_n$ be the eigenvalues of A . Then, the eigenvalues of $I - \alpha A$ are

$$1 - \alpha \lambda_i,$$

and the norm of $I - \alpha A$ is

$$\max_i |1 - \alpha \lambda_i| = |\max(1 - \alpha \lambda_1, 1 - \alpha \lambda_n)|.$$

This is minimized by taking

$$\alpha = \frac{2}{\lambda_n + \lambda_1},$$

in which case the smallest and largest eigenvalues of $I - \alpha A$ become

$$\pm \frac{\lambda_n - \lambda_1}{\lambda_n + \lambda_1},$$

and the norm of $I - \alpha A$ becomes

$$1 - \frac{2\lambda_1}{\lambda_n + \lambda_1}.$$

To show that $\mathbf{x}^t$ converges to the solution, $\mathbf{x}$, consider $\mathbf{x} - \mathbf{x}^t$. We have

$$\begin{aligned}\mathbf{x} - \mathbf{x}^t &= ((I - \alpha A)\mathbf{x} + \alpha \mathbf{b}) - ((I - \alpha A)\mathbf{x}^{t-1} + \alpha \mathbf{b}) \\ &= (I - \alpha A)(\mathbf{x} - \mathbf{x}^{t-1}).\end{aligned}$$

So,

$$\mathbf{x} - \mathbf{x}^t = (I - \alpha A)^t \mathbf{x},$$

and

$$\begin{aligned}\|\mathbf{x} - \mathbf{x}^t\| &= \|(I - \alpha A)^t \mathbf{x}\| \leq \|(I - \alpha A)^t\| \|\mathbf{x}\| \\ &= \|(I - \alpha A)\|^t \|\mathbf{x}\| \\ &\leq \left(1 - \frac{2\lambda_1}{\lambda_n + \lambda_1}\right)^t \|\mathbf{x}\|. \\ &\leq e^{-2\lambda_1 t / (\lambda_n + \lambda_1)} \|\mathbf{x}\|.\end{aligned}$$

So, if we want to get a solution $\mathbf{x}^t$ with

$$\frac{\|\mathbf{x} - \mathbf{x}^t\|}{\|\mathbf{x}\|} \leq \epsilon,$$

it suffices to run for

$$\frac{\lambda_n + \lambda_1}{2\lambda_1} \ln(1/\epsilon) = \left(\frac{\lambda_n}{2\lambda_1} + \frac{1}{2} \right) \ln(1/\epsilon).$$

iterations. The term

$$\frac{\lambda_n}{\lambda_1}$$

is called the *condition number*¹ of the matrix A , when A is symmetric. It is often written $\kappa(A)$, and the running time of iterative algorithms is often stated in terms of this quantity. We see that if the condition number is small, then this algorithm quickly provides an approximate solution.

15.4 A polynomial approximation of the inverse

I am now going to give another interpretation of Richardson's iteration. It provides us with a polynomial in A that approximates A^{-1} . In particular, the t th iterate, $\mathbf{x}^t$ can be expressed in the form

$$p^t(A)\mathbf{b},$$

where p^t is a polynomial of degree $t - 1$. To see this note that

$$\mathbf{x}^t = \mathbf{x} - (I - \alpha A)^t \mathbf{x} = (I - (I - \alpha A)^t) \mathbf{x} = (I - (I - \alpha A)^t) A^{-1} \mathbf{b},$$

so

$$p^t(A) = (I - (I - \alpha A)^t) A^{-1}.$$

This really is a polynomial in A . The term A^{-1} disappears, as the difference

$$(I - (I - \alpha A)^t)$$

is a sum of powers of A of degree at least 1 (the I terms cancel).

We will view $p^t(A)$ as a good approximation of A^{-1} if

$$\|Ap^t(A) - I\|$$

is small. As A , A^{-1} and I all commute, we have

$$Ap^t(A) - I = p^t(A)A - I = (I - (I - \alpha A)^t) - I = (I - \alpha A)^t.$$

Thus, the norm of this matrix is at most

$$\left(1 - \frac{2\lambda_1}{\lambda_n + \lambda_1} \right)^t.$$

¹For general matrices, the condition number is defined to be the ratio of the largest to smallest singular value.

15.5 Faster iterations

It is natural to ask if we can find a faster iterative method, say by exploiting the values of more iterates. To do so, we again use an equation in $\mathbf{x}$ to define an iteration of which $\mathbf{x}$ is a fixed point. For the rest of this lecture, define

$$M \stackrel{\text{def}}{=} I - \alpha A,$$

and remember that

$$\mathbf{x} = M\mathbf{x} + \alpha\mathbf{b}.$$

From this we derive

$$\begin{aligned}\omega\mathbf{x} &= \omega M\mathbf{x} + \omega\alpha\mathbf{b}, \quad \implies \\ \mathbf{x} &= \omega M\mathbf{x} + (1 - \omega)\mathbf{x} + \omega\alpha\mathbf{b}.\end{aligned}$$

This leads us to consider the iteration

$$\mathbf{x}^{t+1} = \omega M\mathbf{x}^t + (1 - \omega)\mathbf{x}^{t-1} + \omega\alpha\mathbf{b}. \quad (15.1)$$

We have chosen this iteration so that $\mathbf{x}$ is a fixed point. Now, let's see how quickly it converges. As before, we find that

$$\mathbf{x} - \mathbf{x}^{t+1} = \omega M(\mathbf{x} - \mathbf{x}^t) + (1 - \omega)(\mathbf{x} - \mathbf{x}^t).$$

So, let's set $\mathbf{y}^t$ to be the t th error vector

$$\mathbf{y}^t = \mathbf{x} - \mathbf{x}^t,$$

and observe that these satisfy

$$\mathbf{y}^{t+1} = \omega M\mathbf{y}^t + (1 - \omega)\mathbf{y}^{t-1}.$$

To express this rule as multiplication by a matrix, we need to consider $\mathbf{y}^{t+1}$ and $\mathbf{y}^t$ together as one vector. We obtain

$$\begin{pmatrix} \mathbf{y}^{t+1} \\ \mathbf{y}^t \end{pmatrix} = \begin{bmatrix} \omega M & (1 - \omega)I \\ I & 0 \end{bmatrix} \begin{pmatrix} \mathbf{y}^t \\ \mathbf{y}^{t-1} \end{pmatrix}.$$

Define

$$S \stackrel{\text{def}}{=} \begin{bmatrix} \omega M & (1 - \omega)I \\ I & 0 \end{bmatrix},$$

and let's find the value of ω that minimizes the absolute values of the eigenvalues of S .

If $\begin{pmatrix} \mathbf{u} \\ \mathbf{v} \end{pmatrix}$ is an eigenvector of S of eigenvalue λ , then

$$\begin{bmatrix} \omega M & (1 - \omega)I \\ I & 0 \end{bmatrix} \begin{pmatrix} \mathbf{u} \\ \mathbf{v} \end{pmatrix} = \begin{pmatrix} \omega M\mathbf{u} + (1 - \omega)\mathbf{v} \\ \mathbf{u} \end{pmatrix} = \begin{pmatrix} \lambda\mathbf{u} \\ \lambda\mathbf{v} \end{pmatrix},$$

so it must be the case that $\mathbf{u} = \lambda\mathbf{v}$. From the top row, we obtain

$$\lambda^2\mathbf{v} = \omega\lambda M\mathbf{v} + (1 - \omega)\mathbf{v}.$$

So, if $\mathbf{v}$ is an eigenvector of M with eigenvalue μ_i , λ will be an eigenvalue of S provided that

$$\lambda^2 = \omega\lambda\mu_i + (1 - \omega) \iff \lambda^2 - \omega\mu_i\lambda + (\omega - 1) = 0.$$

Generically each eigenvalue of M leads to two eigenvalues of S , and by counting we see that this provides all the eigenvalues of S . Solving this quadratic equation, we find eigenvalues

$$\lambda^+ = \frac{\omega\mu_i + \sqrt{\omega^2\mu_i^2 - 4(\omega - 1)}}{2} \quad \text{and} \quad \lambda^- = \frac{\omega\mu_i - \sqrt{\omega^2\mu_i^2 - 4(\omega - 1)}}{2}.$$

Assume that all eigenvalues of M lie strictly between $-\mu$ and μ . We then find that a good choice for ω is

$$\omega = \frac{2}{1 + \sqrt{1 - \mu^2}}.$$

The motivation for this choice is that it causes

$$\omega^2\mu^2 = 4(\omega - 1),$$

so for all μ_i between $-\mu$ and μ

$$\omega^2\mu_i^2 \leq 4(\omega - 1),$$

and the square of the absolute value of the resulting complex eigenvalues λ is

$$|\lambda|^2 = \frac{\omega^2\mu_i^2 + 4(\omega - 1) - \omega^2\mu_i^2}{4} = (\omega - 1).$$

So, all of the eigenvalues of S are complex, and have absolute value

$$\sqrt{(\omega - 1)} = \omega\mu/2.$$

To see that this is an improvement, consider the case in which μ is near 1, say $\mu = 1 - \epsilon$. Then, to first order,

$$\frac{\omega\mu}{2} = \frac{\mu}{1 + \sqrt{1 - \mu^2}} = \frac{1 - \epsilon}{1 + \sqrt{1 - (1 - \epsilon)^2}} \sim \frac{1 - \epsilon}{1 + \sqrt{2\epsilon}} \sim 1 - \epsilon - \sqrt{2\epsilon}.$$

This is much further from 1 than $1 - \epsilon$, and suggests that our algorithm should take around $\sqrt{\kappa(A)}$ iterations to converge, instead of $\kappa(A)$. However, the matrix S was not symmetric, so this argument is not quite precise.

To make it precise, we need to expand our initial vector in the eigenbasis of S . Let $\mathbf{v}_1, \dots, \mathbf{v}_n$ be a basis of eigenvectors of M , and for each let $\mathbf{v}_i^+$ and $\mathbf{v}_i^-$ be the two induced eigenvectors of S :

$$\mathbf{v}_i^+ = \begin{pmatrix} \lambda^+ \mathbf{v}_i \\ \mathbf{v}_i \end{pmatrix} \quad \text{and} \quad \mathbf{v}_i^- = \begin{pmatrix} \lambda^- \mathbf{v}_i \\ \mathbf{v}_i \end{pmatrix}.$$

We will initialize the algorithm by setting

$$\mathbf{x}^0 = \mathbf{0} \quad \text{and} \quad \mathbf{x}^1 = \mathbf{b}.$$

15.6 A polynomial from the second-order method

To make this precise, we will see that this method also applies a polynomial in A to b , and that it is a very good approximation of A^{-1} . Our analysis will be facilitated by a judicious choice of $\mathbf{x}^0$ and $\mathbf{x}^1$. We set

$$\mathbf{x}^0 = \mathbf{0} \quad \text{and} \quad \mathbf{x}^1 = \mathbf{b}.$$

We then find that

$$\mathbf{y}^i = \begin{bmatrix} I & 0 \end{bmatrix} S^{i-1} \begin{bmatrix} M \\ I \end{bmatrix} \mathbf{y}^0.$$

So, there is a polynomial p^i of degree i such that $\mathbf{y}^i = p^i(A)\mathbf{x}$. To bound the eigenvalues of $p^i(A)$. With some computation, one can prove that the eigenvalues of $p^i(A)$ are real and have absolute value at most

$$\frac{\mu\omega}{2}(1 + i\sqrt{1 - \mu^2}).$$

I thought I had a way of making this computation obvious, but I was wrong. Sorry about that.

15.7 The best polynomials

To evaluate how close $p(A)$ is to A^{-1} , we measured the eigenvalues of

$$I - Ap(A).$$

If we want to find the polynomial in A of degree $i - 1$ that comes closest to A^{-1} , we can instead set

$$q(x) = 1 - xp(x),$$

and find the polynomial q that minimizes

$$|q(\lambda_i)|$$

for λ_i and eigenvalue of A . Of course, we need the constant $q(x)$ to be 1, so we restrict $q(0) = 1$. It is possible to find the optimal polynomial q under the assumption that the eigenvalues of A lie between λ_1 and λ_n . If $\lambda_1 > 0$, then the optimal polynomial of degree i may be constructed using Chebyshev polynomials, and guarantees that

$$|q^i(\lambda)| \leq 2 \left(1 - \frac{2}{\sqrt{\lambda_n/\lambda_1} + 1} \right)^i,$$

for all $\lambda \in [\lambda_1, \lambda_n]$. So, to get error ϵ , the best polynomial will require degree around

$$\ln(1/\epsilon) \left(\sqrt{\lambda_n/\lambda_1} + 1 \right) / 2.$$

This is quadratically better than we got from the Richardson iteration.

15.8 The Conjugate Gradient

For positive-definite systems, there is an algorithm that dominates all of these: the Conjugate Gradient. I do not have time to explain how it works, but I can tell you what it does. In the i th iteration, it finds the optimal approximation $\mathbf{x}^i$ to $\mathbf{x}$ in the span of

$$\{\mathbf{b}, A\mathbf{b}, A^2\mathbf{b}, \dots, A^{i-1}\mathbf{b}\}.$$

It does this while performing an amount of work comparable to that done by the second-order method: one matrix multiplication and a couple of vector operations per iteration.

However, the Conjugate Gradient defines optimality in a slightly counter-intuitive way. It minimizes

$$\|\mathbf{x}^i - \mathbf{x}\|^A \stackrel{\text{def}}{=} \sqrt{(\mathbf{x}^i - \mathbf{x})^T A (\mathbf{x}^i - \mathbf{x})}.$$

However, for many applications this is actually a better way to measure error than the ordinary Euclidean norm!

By using Chebyshev polynomials, we can show that the Conjugate Gradient method will always find a vector $\mathbf{x}^i$ in the i th iteration that satisfies

$$\|\mathbf{x} - \mathbf{x}^i\|_A \leq 2 \left(1 - \frac{2}{\sqrt{\lambda_n/\lambda_1} + 1} \right)^i \|\mathbf{x}\|_A.$$

15.9 Solving equations in Laplacians

It may seem strange to solve linear equations in Laplacian matrices, since they are degenerate. It becomes less strange once you observe that we know what their nullspace is, and can restrict ourselves to working orthogonal to this nullspace. In this case, we only care about the non-zero eigenvalues. The problem of solving linear equations in Laplacians arises in many applications. I hope I have time to mention some in class.

In the meantime, let's look at what happens when we try to solve an equation in the Laplacian of a path graph. If we are going to use an iterative method, it is clear that we will need at least $\Theta(n)$ iterations. After all, it takes $n/2$ iterations before the value of $\mathbf{x}(1)$ has any dependence on $\mathbf{b}(n/2)$. On the other hand, the unweighted path graph has $\lambda_n/\lambda_2 = O(n^2)$, so an iterative method can converge in $O(n)$ iterations. This tells us that, up to constant factors, the dependence of the convergence rate on λ_n/λ_2 cannot be improved.

On the other hand, we see that if λ_2 is large, then we can solve equations in the Laplacian very quickly.

In the next lecture, I will show you a technique for avoiding a dependence on λ_n/λ_2 entirely. By approximating Laplacians by Laplacians of simpler graphs, we will see how to solve any Laplacian system in time $O(m^{4/3})$.

15.10 Research Project

The reason that I taught this lecture this way is that I have been trying to find a more intuitive explanation for the quadratic speedup of the Chebyshev method over Richardson's. I don't like merely appealing to the properties of Chebyshev polynomials. In fact, I would really prefer a geometric explanation. Please help me find one!