

Convergence in $\mathbb{R}^n$

We say a sequence of vectors $v^{(k)} \in \mathbb{R}^n$ converges to $v \in \mathbb{R}^n$ as $k \rightarrow \infty$ if

$$\lim_{k \rightarrow \infty} \|v^{(k)} - v\| = 0$$

In $\mathbb{R}^n$ the choice of norm does not matter as all norms are equivalent: (this is true for any finite dimensional space)

Two norms $\|\cdot\|$ and $\|\cdot\|'$ are equivalent if there are two constants c_1 and $c_2 > 0$ s.t.

$$c_1 \|x\| \leq \|x\|' \leq c_2 \|x\|$$

Also $\mathbb{R}^n$ is complete meaning any sequence satisfying the Cauchy criterion converges.

$$\forall \epsilon > 0 \exists N \text{ s.t. } \forall i, j \geq N \quad \|v^{(i)} - v^{(j)}\| < \epsilon$$

Neumann Series

Let $A \in \mathbb{R}^{n \times n}$ be such that $\|A\| < 1$, then

i) $I - A$ is invertible

$$\text{ii) } (I - A)^{-1} = \sum_{k=0}^{\infty} A^k$$

Proof: Assume for contradiction that $I - A$ is singular

$$\Rightarrow \exists z \in \mathbb{R}^n \text{ s.t. } (I - A)z = 0, \text{ take w.l.o.g. } z \text{ s.t. } \|z\| = 1$$

then:

$$1 = \|z\| = \|Az\| \leq \overbrace{\|A\|}^{< 1} \|z\| < \|z\| = 1. \text{ Contradiction!}$$

Now we need to show that

$$\sum_{k=0}^m A^k \rightarrow (I-A)^{-1} \quad (\text{ie. partial sums converge})$$

Notice:

$$(I-A) \sum_{k=0}^m A^k = \sum_{k=0}^m A^k - A^{k+1} = A^0 - A^{m+1} = I - A^{m+1}$$

↙ telescoping series.

Now using properties of induced matrix norm: ($\|AB\| \leq \|A\| \|B\|$)

$$\|A^{m+1}\| \leq \|A\|^m \rightarrow 0$$

Thus:

$$\left\| (I-A) \sum_{k=0}^m A^k - I \right\| \leq \|A\|^{m+1} \rightarrow 0 \text{ as } m \rightarrow \infty.$$

Q.E.D.

Note: This theorem is a very intuitive generalization of what you already know about geometric series.

$$\frac{1}{1-x} = 1 + x + x^2 + \dots, \text{ when } |x| < 1.$$

Here is the kind of result we can show with Neumann Series:

$$\|(I-A)^{-1}\| \leq \sum_{k=0}^{\infty} \|A^k\| \leq \sum_{k=0}^{\infty} \|A\|^k = \frac{1}{1-\|A\|}.$$

Note: This theorem is very useful for the theory, however it is seldom used in practice. (There are faster iterative methods for solving $Ax = b$).

Here is an easy but handy generalization of

Neuman Series:

Theorem

If A and B are matrices s.t. $\|I - AB\| < 1$ then
 A, B are invertible and

$$A^{-1} = B \sum_{k=0}^{\infty} (I - AB)^k$$

$$B^{-1} = \left[\sum_{k=0}^{\infty} (I - AB)^k \right] A$$

Proof: By Neumann Series:

$I - (I - AB) = AB$ is invertible and:

$$(AB)^{-1} = \sum_{k=0}^{\infty} (I - AB)^k$$

$$\Rightarrow A^{-1} = B B^{-1} A^{-1} = B \sum_{k=0}^{\infty} (I - AB)^k$$

$$B^{-1} = B^{-1} A^{-1} A = \left[\sum_{k=0}^{\infty} (I - AB)^k \right] A$$

Iterative refinement

(30)

4

Let $x^{(0)}$ be an approx. sol to

$$Ax = b.$$

then:

$$x = x^{(0)} + A^{-1} (b - Ax^{(0)})$$

We often call

$$r^{(0)} = b - Ax^{(0)} = \text{residual vector}$$

$$e^{(0)} = A^{-1} r^{(0)} = \text{error vector.}$$

Algorithm

$x^{(0)}$ given

for $k = 0, 1, \dots$

$$\begin{cases} r^{(k)} = b - Ax^{(k)} \\ Ae^{(k)} = r^{(k)} \\ x^{(k+1)} = x^{(k)} + e^{(k)} \end{cases}$$

————— usually inexact solve

Solution can be improved even if the solves are inexact for example:

→ LU factorization is stopped prematurely ($\equiv$ incomplete LU factor)

→ $Ae^{(k)} = r^{(k)}$ is solved with an iterative method

→ A is not known exactly.

To be more precise about what we mean by "solving approximately"

Let $B \equiv$ "approximate" inverse of A

$$\approx A^{-1} \text{ (in some sense we shall see)}$$

The iterates are:

$$x^{(0)} = Bb$$

$$x^{(k+1)} = x^{(k)} + B(b - Ax^{(k)}), \quad k \geq 0.$$

Looking at a few iterates we see a pattern emerge.

$$x^{(0)} = Bb$$

$$x^{(1)} = x^{(0)} + B(b - Ax^{(0)})$$

$$= Bb + B(b - ABb)$$

$$= Bb + B(I - AB)b$$

$$x^{(2)} = x^{(1)} + B(b - Ax^{(1)})$$

$$= Bb + B(I - AB)b + B(b - A(Bb + B(I - AB)b))$$

$$= Bb + B(I - AB)b + B((I - AB)b - AB(I - AB)b)$$

$$= Bb + B(I - AB)b + B(I - AB)^2 b$$

$$\boxed{x^{(m)} = B \sum_{k=0}^m (I - AB)^k b}$$

Theorem (Iterative refinement)

Iterates from iterative refinement are:

$$x^{(m)} = B \sum_{k=0}^m (I - AB)^k b \quad (m \geq 0)$$

and because of the generalized Neuman series:

$$x^{(m)} \rightarrow x \quad \text{as } m \rightarrow \infty \quad \text{when } \|I - AB\| < 1.$$

Proof: Can be done by induction on m .

Solving equations with iterative methods

(32)
6

Idea: Instead of having a direct method which finds solution in finitely many steps, use an iterative method to find successive approx that converge to sol.

Motivational example

$$\begin{bmatrix} 7 & -6 \\ -8 & 9 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} 3 \\ -4 \end{bmatrix}$$

Jacobi method: solve i -th equation for i -th unknown:

$$\begin{cases} x_1^{(k)} = \frac{3}{7} + \frac{6}{7} x_2^{(k-1)} \\ x_2^{(k)} = \frac{-4}{9} + \frac{8}{9} x_1^{(k-1)} \end{cases}$$

Gauss Seidel method: use improved value of $x_1^{(k)}$ in second equation.

$$\begin{cases} x_1^{(k)} = \frac{3}{7} + \frac{6}{7} x_2^{(k-1)} \\ x_2^{(k)} = \frac{-4}{9} + \frac{8}{9} x_1^{(k)} \end{cases} \quad \text{converges faster!}$$

Note: convergence for these methods depends on initial guess $(x_1^{(0)}, x_2^{(0)})$

Matrix splitting

Consider the system $Ax = b$.

Introduce a non-singular matrix Q and rewrite system as:

$$Qx = (Q - A)x + b$$

Get iteration:

$$Qx^{(k)} = (Q - A)x^{(k-1)} + b, \quad k \geq 1, \quad x^{(0)} \text{ given}$$

Q should have the desirable properties:

- ① $x^{(k)} \rightarrow x$ where x is a sol. of $Ax=b$
- ② Solving systems with Q is cheap.

If $x^{(k)} \rightarrow x$ then the limit x must satisfy:

$$Qx = (Q-A)x + b \Rightarrow Ax=b$$

Assuming A & Q are non-singular:

$$x^{(k)} = (I - Q^{-1}A)x^{(k-1)} + Q^{-1}b$$

$$x = (I - Q^{-1}A)x + Q^{-1}b$$

$$x^{(k)} - x = (I - Q^{-1}A)(x^{(k-1)} - x)$$

$$\begin{aligned} \Rightarrow \|x^{(k)} - x\| &\leq \|I - Q^{-1}A\| \|x^{(k-1)} - x\| \\ &\vdots \\ &\leq \|I - Q^{-1}A\|^k \|x^{(0)} - x\| \end{aligned}$$

Thus when $\|I - Q^{-1}A\| < 1$ we have $x^{(k)} \rightarrow x$, where $Ax=b$.

Theorem If $\|I - Q^{-1}A\| < 1$ for some induced matrix norm then seq. $Qx^{(k)} = (Q-A)x^{(k-1)} + b$

converges to sol. of $Ax=b$ for any starting vector $x^{(0)}$.

Richardson method

(30)

8

$$Q = I$$

$$x^{(k)} = (I - A)x^{(k-1)} + b = x^{(k-1)} + \underbrace{b - Ax^{(k-1)}}_{r^{(k-1)}}$$

converges when $\|I - A\| < 1$ in some induced matrix norm.

Jacobi method

$$Q = \begin{bmatrix} a_{11} & & & & \\ & a_{22} & & & \\ & & \ddots & & \\ & & & a_{nn} & \end{bmatrix} = \text{diag}(\text{diag}(A)) = \text{diagonal matrix with same entries as diagonal of } A.$$

$$Q^{-1}A = \begin{bmatrix} 1 & a_{12}/a_{11} & a_{13}/a_{11} & \dots & a_{1n}/a_{11} \\ a_{21}/a_{22} & 1 & a_{23}/a_{22} & \dots & a_{2n}/a_{22} \\ \vdots & & & & \\ & & & a_{n-1,n}/a_{n-1,n-1} & \\ a_{n1}/a_{nn} & \dots & \dots & \dots & 1 \end{bmatrix}$$

Thus:

$$\|I - Q^{-1}A\|_{\infty} = \max_{1 \leq i \leq n} \sum_{\substack{j=1 \\ j \neq i}}^n \frac{|a_{ij}|}{|a_{ii}|}.$$

Theorem If A is diagonally dominant, then the Jacobi method's iterates converge to a solution to $Ax = b$.

Proof.

$$A \text{ diag dominant} \Rightarrow |a_{ii}| > \sum_{\substack{j=1 \\ j \neq i}}^n |a_{ij}|, \quad 1 \leq i \leq n$$

$$\Rightarrow \|I - Q^{-1}A\|_{\infty} < 1$$

$\Rightarrow$ convergence.

Linear algebra review

(35)

9

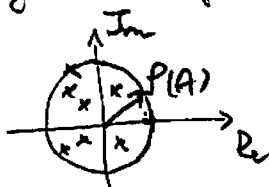
Eigenvalues of A = roots of characteristic polyn. of A

The characteristic poly. of A is defined by:

$$P(\lambda) = \det(\lambda I - A)$$

Spectral radius:

$$\begin{aligned} \rho(A) &= \max \{ |\lambda| \mid \det(\lambda I - A) = 0 \} \\ &= \text{radius of smallest circle in } \mathbb{C} \text{ containing} \\ &\quad \text{all eigenvalues of } A. \end{aligned}$$



Two matrices A and B are similar if there is an invertible matrix X s.t.

$$AX = XB.$$

note: Two similar matrices have same eigenvalues:

$$\begin{aligned} P_A(\lambda) &= \det(\lambda I - A) = \det(X^{-1}(\lambda I - A)X) \\ &= \det(\lambda I - \underbrace{X^{-1}AX}_B) = P_B(\lambda) \end{aligned}$$

Here we used the facts:

$$\det(AB) = \det(A) \det(B)$$

$$\det(A^{-1}) = \det(A)^{-1}.$$

Before we continue studying iterative methods for solving linear systems we need two technical results.

Theorem Every matrix is similar to an upper triangular matrix with arbitrarily small off-diagonal elements.

Proof: Use Schur factorization

$$A = Q T Q^T, \text{ where } Q^T Q = I$$

this is an eigenvalue revealing factorization as the eigenvalues of A appear on diagonal of T . we will see more about this factorization when we study eigenvalue algorithms.

$$\text{Let } D = \begin{bmatrix} \varepsilon & & \\ & \varepsilon^2 & \\ & & \ddots \\ & & & \varepsilon^n \end{bmatrix}$$

$$\text{Then } (D^{-1} T D)_{ij} = t_{ij} \varepsilon^{j-i}$$

But $T = (\tau)$ thus:

$$t_{ij} = 0 \text{ when } i > j \text{ (elements below diag are zero)}$$

Now for elements above diag ($i < j$):

$$|t_{ij} \varepsilon^{\tilde{j-i}}| \leq \varepsilon |t_{ij}|$$

Thus

$$A = \underbrace{Q D}_X \underbrace{D^{-1} T D}_B \underbrace{D^{-1} Q^T}_{X^{-1}}. \quad \text{QED}$$

The second technical result is:

37
11

Theorem

$$P(A) = \inf_{\|\cdot\|} \|A\|, \text{ where } \inf \text{ is over all induced matrix norms.}$$

This means $P(A) = \|A\|_2$ = "smallest" induced norm one can choose to measure A .

proof:

$$\bullet \underline{P(A) \leq \inf_{\|\cdot\|} \|A\|}$$

Let $\|\cdot\|$ be a vector norm and let λ, x be an eigenvalue/vector of A , then:

$$\|A\| \|x\| \geq \|Ax\| = \|\lambda x\| = |\lambda| \|x\|$$

$$\Rightarrow |\lambda| \leq \|A\|. \quad (\text{since eigenvector } x \neq 0)$$

Since this is true for any eigenvalue λ we must have:

$$P(A) \leq \|A\|$$

Since this is true for any induced matrix norm $\|\cdot\|$, it must be true for the least upper bound as well:

$$P(A) \leq \inf_{\|\cdot\|} \|A\|.$$

$$\bullet \underline{P(A) \geq \inf_{\|\cdot\|} \|A\|}$$

: $\forall \epsilon > 0$, there is a non-singular matrix S s.t.

$$S^{-1}AS = D + T, \text{ where } \|T\|_\infty \leq \epsilon$$

(\) (\)
 (strict)

$$\Rightarrow \|S^{-1}AS\|_\infty = \|D + T\|_\infty \leq \underbrace{\|D\|_\infty}_{=P(A)} + \underbrace{\|T\|_\infty}_{\leq \epsilon} \leq P(A) + \epsilon$$

To conclude we notice (this is easily shown) that:

$$\|A\|'_\infty = \|S^{-1}AS\|_\infty \text{ is an induced matrix norm}$$

$$\Rightarrow \inf_{\|\cdot\|} \|A\| \leq \|A\|'_\infty \leq \rho(A) + \epsilon$$

$$\Rightarrow \inf_{\|\cdot\|} \|A\| \leq \rho(A) \text{ since inequality above holds } \forall \epsilon > 0.$$

QED

We are now ready for the main result of this section. If you notice all convergence results we have given up to now are simply sufficient conditions for convergence. Here is necessary and sufficient condition for convergence.

First let us write the iterative method in a slightly more general way:

$$x^{(k)} = Gx^{(k-1)} + c \quad (*)$$

The matrix splitting methods (Jacobi etc.) we have seen before fall on this framework:

$$Qx^{(k)} = (Q - A)x^{(k-1)} + b$$

$$x^{(k)} = \underbrace{(I - Q^{-1}A)}_{= G} x^{(k-1)} + \underbrace{Q^{-1}b}_{= c}$$

What is limit of iteration (*)?

$$x = Gx + c \Rightarrow (I - G)x = c$$

$$\boxed{x = (I - G)^{-1}c}$$

Theorem The iterative method

$$x^{(k)} = G x^{(k-1)} + c$$

converges to $(I - G)^{-1}c$ if and only if $\rho(G) < 1$. essentially "sharpest" convergence measure

proof:

- Assume $\rho(G) < 1$.

Since $\rho(G) = \inf_{\|\cdot\|} \|G\|$, there is an induced matrix norm for which $\|G\| < 1$.

$$x^{(1)} = G x^{(0)} + c$$

$$x^{(2)} = G^2 x^{(0)} + Gc + c$$

$$x^{(3)} = G^3 x^{(0)} + G^2 c + Gc + c$$

$\vdots$

$$x^{(k)} = G^k x^{(0)} + \sum_{j=0}^{k-1} G^j c$$

Note:

$$\|G^k x^{(0)}\| \leq \|G\|^k \|x^{(0)}\| \rightarrow 0 \text{ as } k \rightarrow \infty$$

and

$$\sum_{j=0}^{\infty} G^j c = (I - G)^{-1} c \text{ by Neuman series}$$

Hence $x^{(k)} \rightarrow (I - G)^{-1} c$, regardless of initial guess $x^{(0)}$.

- Assume $\rho(G) \geq 1$: Take $u \neq 0$ s.t.

$$Gu = \lambda u \text{ and } |\lambda| \geq 1.$$

Then with $x^{(0)} = 0$ and $c = u$:

then:
$$z^{(k)} = \sum_{j=0}^{k-1} G^j u = \sum_{j=0}^{k-1} \lambda^j u = \begin{cases} \pm k u & \text{if } \lambda = \pm 1 \\ \frac{1 - \lambda^k}{1 - \lambda} & \text{if } |\lambda| \neq 1 \end{cases}$$

(40)
14

then $z^{(k)}$ does not converge as $k \rightarrow \infty$. QED.

Corollary The matrix splitting iteration:

$$Q z^{(k)} = (Q - A) z^{(k-1)} + b$$

converges to solution of $Ax = b$

if and only if

$$\rho(I - Q^{-1}A) < 1.$$

Gauss Seidel method:

Q = lower triangular part of A (including diag)
= easy to solve systems (forward substitution)

Theorem If A is diag. dominant then Gauss-Seidel method converges to sol. of $Ax = b$, for any starting vector $x^{(0)}$.

Proof: we need to show $\rho(I - Q^{-1}A) < 1$

Let λ be an eigenvalue of $I - Q^{-1}A$ and x a corresp. eigenvector chosen s.t. $\|x\|_\infty = 1$.

$$(I - Q^{-1}A)x = \lambda x \Leftrightarrow \overbrace{(Q - A)x}^{\text{strictly upper triangular part}} = \lambda Qx$$

$$\Leftrightarrow - \sum_{j=i+1}^n a_{ij} x_j = \lambda \sum_{j=1}^i a_{ij} x_j$$

for $i = 1, \dots, n$

$$\Rightarrow \lambda a_{ii} x_i = - \lambda \sum_{j=1}^{i-1} a_{ij} x_j - \sum_{j=i+1}^m a_{ij} x_j \quad (i=1, \dots, n) \quad \textcircled{41} \quad 15$$

Let i be s.t. $|x_i|=1$ (i.e. element realizing $\max_i |x_i|=1$)

$$|\lambda| |a_{ii}| \leq |\lambda| \sum_{j=1}^{i-1} |a_{ij}| + \sum_{j=i+1}^m |a_{ij}|$$

$$\Rightarrow |\lambda| \leq \frac{\sum_{j=i+1}^m |a_{ij}|}{|a_{ii}| - \sum_{j=1}^{i-1} |a_{ij}|} < 1 \quad \text{by diagonal dominance.}$$

$$\Rightarrow \rho(I - Q^{-1}A) < 1 \Rightarrow \text{convergence.} \quad \text{QED.}$$

$$\text{Let } A = D - E - F$$

$$\underbrace{(\diagup)}_{\text{strict}} \underbrace{(D)}_{\text{strict}} \underbrace{(\diagdown)}$$

Successive over relaxation (SOR):

$$Q = \frac{1}{\omega} (D - \omega E), \text{ where } \omega > 0.$$

Iteration becomes:

$$\begin{aligned} \frac{1}{\omega} (D - \omega E) x^{(k)} &= \left(\frac{1}{\omega} (D - \omega E) - D + E + F \right) x^{(k)} + b \\ &= \frac{1}{\omega} \left((1 - \omega) D + \omega F \right) x^{(k)} + b \end{aligned}$$

$$(D - \omega E) x^{(k)} = ((1 - \omega) D + \omega F) x^{(k)} + \omega b$$

Complexity is similar to Gauss-Seidel, but because of extra parameter ω one can get better convergence or even convergence in some systems where G-S does not converge.

Note: When $\omega=1$ SOR $\equiv$ Gauss Seidel.

(42)

16

Here is another way of looking at SOR:

$$(1-\omega)x \left[D x^{(k)} = D x^{(k-1)} \right] \quad \text{identity.}$$

$$\omega x \left[(D-E) x^{(k)} = F x^{(k-1)} + b \right] \quad \text{G-S.}$$

$$(D - \omega E) x^{(k)} = ((1-\omega)D + \omega F) x^{(k-1)} + \omega b$$

Depending on matrix $A = D - E - F$, there are ways of choosing ω optimally.

Extrapolation methods:

$$\gamma x \left[x^{(k)} = G x^{(k-1)} + c \right]$$

$$(1-\gamma)x \left[x^{(k)} = x^{(k-1)} \right]$$

$$\begin{aligned} x^{(k)} &= \gamma (G x^{(k-1)} + c) + (1-\gamma) x^{(k-1)} \\ &= G_\gamma x^{(k-1)} + \gamma c \end{aligned}$$

where $G_\gamma = \gamma G + (1-\gamma)I$
= interp between G and I .

If iteration converges then:

$$x = \gamma (Gx + c) + (1-\gamma)x$$

$$\Rightarrow x = Gx + c$$

So setting $G = (I - Q^{-1}A)$ and $c = Q^{-1}b$ we can solve $Ax = b$.

Question How do we choose parameter γ in

$$G_\gamma = \gamma G + (1-\gamma)I$$

$$\text{s.t. } P(G_\gamma) = \text{minimum}$$

(the smaller $P(G_\gamma)$ is the faster convergence rate)

If we know

$$a \leq \lambda(G) \leq b$$

why?

then

$$\lambda(G_\gamma) \text{ is between } \gamma a + (1-\gamma)1 \text{ and } \gamma b + (1-\gamma)1.$$

$$\begin{aligned} \Rightarrow P(G_\gamma) &= \max |\lambda(G_\gamma)| \quad \checkmark \text{ check.} \\ &= \max |\gamma \lambda(G) + 1 - \gamma| \\ &\leq \max_{a \leq \lambda \leq b} |\gamma \lambda + 1 - \gamma|. \end{aligned}$$

This problem has explicit solution:

Theorem If $\lambda(G) \in [a, b]$, and $1 \notin [a, b]$, then the best choice for γ is:

$$\gamma_* = \frac{2}{2-a-b} \quad \text{and} \quad P(G_\gamma) \leq (1-|\gamma|)d$$

where d = distance from 1 to $[a, b]$.

We shall not prove this result.

However be aware one can think of more complicated functionals of G and ask the question of which one gives faster convergence

Here is one example that has ties to Chebyshev polynomials which you should have seen in Math 5610, when one is allowed to move the interpolation nodes in order to reduce the Lagrange interpolation error.

Chebyshev acceleration idea

We start with iterative method:

$$x^{(k)} = G x^{(k-1)} + c$$

and ask whether a linear combination of previous iterates

$$u^{(k)} = \sum_{i=0}^k a_i^{(k)} x^{(i)}$$

with coeff s.t. $\sum_{i=0}^k a_i^{(k)} = 1$ is a better approx

than the $x^{(0)}, x^{(1)}, \dots, x^{(k)}$.

$$u^{(k)} = \sum_{i=0}^k a_i^{(k)} x^{(i)}$$

$$\ominus \quad x = \sum_{i=0}^k a_i^{(k)} x$$

$$u^{(k)} - x = \sum_{i=0}^k a_i^{(k)} (x^{(i)} - x)$$

$$= \sum_{i=0}^k a_i^{(k)} G^i (x^{(0)} - x)$$

$$= P(G) (x^{(0)} - x) \Rightarrow \|u^{(k)} - x\| \leq \|P(G)\| \|x^{(0)} - x\|$$

where P is the polynomial:

$$P(z) = \sum_{i=0}^k a_i^{(k)} z^i$$

note: $P(A)$ means exactly what you would have expected.

$$P(A) = \sum_{i=0}^k a_i^{(k)} A^i, \text{ with } A^0 \equiv I.$$

Now it is not hard to show that

$$\lambda(A) \equiv \text{eigenvalue of } A \Rightarrow P(\lambda(A)) \equiv \text{eigenvalue of } P(A)$$

Thus if we want best approximation we must have:

$$\begin{aligned} P(P(G)) &= \max_{1 \leq i \leq n} |P(\lambda_i(G))| \\ &\leq \max_{z \in S} |P(z)| \end{aligned}$$

where polynomial P is s.t.

$$P(1) = 1 \quad (\text{since this is equiv. to } \sum_{i=0}^k a_i^{(k)} = 1)$$

$\Rightarrow P$ is a monic polynomial

Morale of story is that we have reduced the problem of finding the best approx in terms of previous iterates to the problem of finding a polynomial with smallest maximum over some set S .

When $S = [-1, 1]$ the solution can be given in terms of Chebyshev polynomials:

$$\begin{aligned} T_k(z) &\equiv \text{unique polynomial minimizing } \max_{z \in [-1, 1]} |T_k(z)| \\ &\text{with leading coefficient being } 2^{k-1}. \end{aligned}$$

$$\equiv \cos(k \cos^{-1} z)$$

why? Because

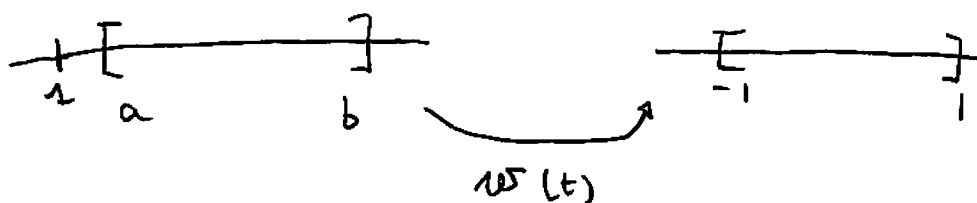
(46)
20

$$T_k(z) = \min_{\substack{P \in \Pi_k \\ \text{largest coeff} \\ = 2^{k-1}}} \max_{z \in [-1, 1]} |P(z)|$$

it follows that for $\beta \in \mathbb{R} \setminus (-1, 1)$: (to avoid dividing by zero below)
(all roots of T_k are in $(-1, 1)$)

$$\frac{T_k(z)}{T_k(\beta)} = \min_{\substack{P \in \Pi_k \\ P(\beta) = 1}} \max_{z \in [-1, 1]} |P(z)|$$

If $S = [a, b]$ = interval where eigenvalues of G lie,
and $[a, b] \not\subset 1$, then we can find transformation s.t.:



where

$$w(t) = (-1) \frac{t-b}{a-b} + (1) \frac{a-t}{a-b} \quad (\text{Lagrange interp})$$
$$= \frac{a+b-2t}{a-b}$$

then:

$$\frac{T_k(w(t))}{T_k(w(1))} = \min_{\substack{P \in \Pi_k \\ P(1) = 1}} \max_{z \in [a, b]} |P(z)|$$

$\hookrightarrow$ can be computed on the fly with a three term recurrence.

So we only need to remember 2 past iterates.

CONJUGATE GRADIENT

44

21

Objective: Solve $Ax = b$ when $A \in \mathbb{R}^{n \times n}$, symm. pos. def.

We shall use the notation $\langle x, y \rangle = x^T y = \sum_{i=1}^n x_i y_i$

- inner product.

Recall properties of an inner product (these hold in general not only in $\mathbb{R}^n$)

inner product $\equiv$ positive symmetric bilinear form

i) $\langle x, y \rangle = \langle y, x \rangle$ (symm)

ii) $\langle x, x \rangle \geq 0$ for all $x \in \mathbb{R}^n$ and $\langle x, x \rangle = 0 \Leftrightarrow x = 0$ } (pos.)

iii) $\langle \alpha x + \beta y, z \rangle = \alpha \langle x, z \rangle + \beta \langle y, z \rangle$ (bilin)

Also we have the following property.

iv)

$$\langle x, Ay \rangle = x^T (Ay) = (A^T x)^T y = \langle A^T x, y \rangle$$

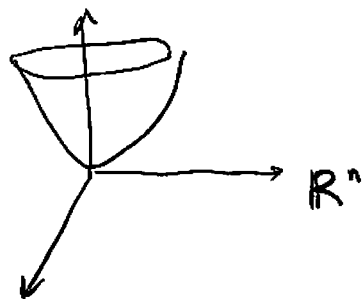
Conjugate gradient exploits the following equivalence between linear system and an optimization problem:

(*) $\boxed{\text{Find } x \text{ s.t. } Ax = b} \Leftrightarrow \boxed{\min_{x \in \mathbb{R}^n} \frac{1}{2} x^T A x - x^T b = q(x)}$

Here $q(x) \equiv$ quadratic function. When $x \in \mathbb{R}^n$:

(48)

22



Finding solutions of $Ax=b \Leftrightarrow$ finding bottom of bowl.

note: the fact that we have an (upward) bowl and not a hyperboloid is due to A being s.p.d.

$(x^T Ax \geq 0 \quad \forall x \in \mathbb{R}^n, \text{ so we cannot have branches going to } -\infty)$

If you know some vector calculus and first and second order sufficient conditions for having a minimizer then showing (*) ... is simply:

$$\nabla q(x) = Ax - b$$

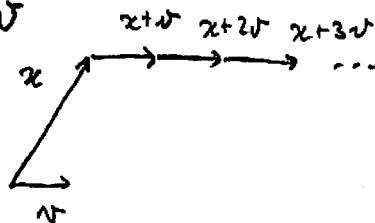
$$\nabla^2 q(x) = A = \text{s.p.d.}$$

Think of $f: \mathbb{R} \rightarrow \mathbb{R}; f \in C^2$
 $f'(x_*) = 0$ and $f''(x_*) > 0$
 $\Rightarrow x_*$ is a minimizer of f

However here is a direct proof of (*)

proof of (*):

Look at ray $x + tv$



$$\begin{aligned}
 q(x+tv) &= \frac{1}{2} (x+tv)^T A (x+tv) - (x+tv)^T b \\
 &= q(x) + t v^T A x - t v^T b + \frac{1}{2} t^2 v^T A v
 \end{aligned}$$

$$\frac{d}{dt} q(x+tv) = v^T (Ax - b) + t v^T A v$$

$$q \text{ is minimum when } t = t_* = - \frac{v^T (Ax - b)}{v^T A v}$$

and:

$$\begin{aligned}
 q(x+t_*v) &= q(x) + t_* (v^T (Ax - b) + \frac{1}{2} t_* v^T A v) \\
 &= q(x) + t_* (v^T (Ax - b) - \frac{1}{2} v^T (Ax - b)) \\
 &= q(x) - \frac{(v^T (Ax - b))^2}{2 v^T A v}
 \end{aligned}$$

What does this mean?

If x is minimizer of $q(x)$, then if we move from x in any direction v , then we should not be able to get a smaller value for q .

$$\Rightarrow v^T (Ax - b) = 0 \quad \text{for all directions } v$$

$$\Rightarrow Ax - b = 0$$

Conversely, If $Ax = b$ then

$$q(x+tv) = q(x) + \frac{1}{2} t^2 v^T A v > q(x) \quad \text{when } v \neq 0.$$

If $x^{(k)}$ is given then our calculation reveals that the best possible stepsize we can take (i.e. the one that reduces the most value of objective function $q(x)$) is:

where $r^{(k)} = b - Ax^{(k)}$ = residual at step k .

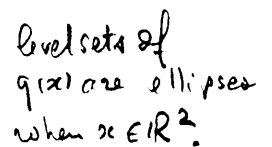
One example of such method is:

Steepest descent where $v^{(k)} = r^{(k)} = -\nabla q(x^{(k)})$

Algorithm.

for $k = 1, 2, \dots$

This is not a very efficient method. It has slow, typically "zig-zagging" convergence to solution:



Note: $r(h)$ is $\perp$ to level sets of $q(x)$.

Conjugate Gradient method (Hestenes & Stiefel, 1952)

(51)
25

$x_{k+1} = x_k + \alpha_k p_k$ where α_k, p_k is determined s.t.
 x_{k+1} solves

$$(1) \quad \min \frac{1}{2} x^T A x - b^T x$$

$$x \in x_0 + \text{span} \{ p_0, \dots, p_{k-1}, r_k \}$$

$$x = \hat{x} + x_0$$

$$\Rightarrow (2) \quad \min \frac{1}{2} \hat{x}^T A \hat{x} - \hat{x}^T b$$

$$\hat{x} \in \text{span} \{ p_0, \dots, p_{k-1}, r_k \}$$

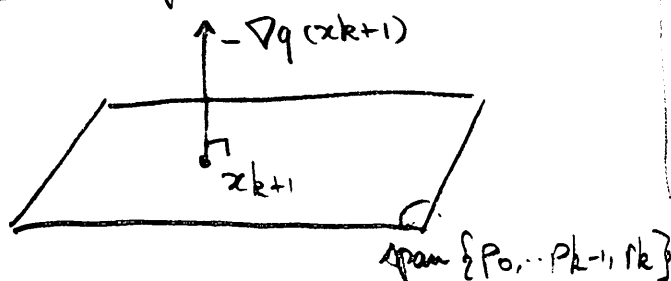
Recall: $\text{span} \{ u_1, u_2, \dots, u_n \} = \left\{ \sum_{i=1}^n \alpha_i u_i \mid \alpha_i \in \mathbb{R} \text{ for } i=1, \dots, n \right\}$
 = set of all possible linear combinations of vectors $u_1, \dots, u_n$
 = linear span of family $\{ u_1, \dots, u_n \}$.

It can be shown that $\hat{x}_{k+1}$ solves (2) iff:

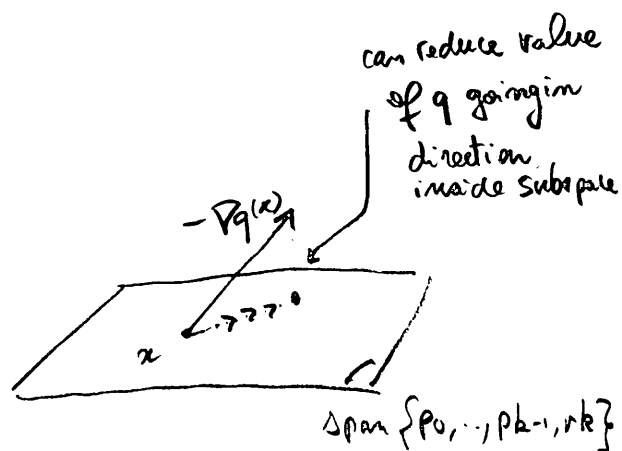
$$(A \hat{x}_{k+1} - b)^T v = 0 \quad \forall v \in \text{span} \{ p_0, \dots, p_{k-1}, r_k \}$$

$$\Rightarrow \underbrace{(A x_{k+1} - b)^T}_{= -\nabla q(x_{k+1})} v = 0$$

Intuitively:



OPTIMAL: no direction of descent in $\text{span} \{ \dots \}$



SUB-OPTIMAL

note: v is called a descent direction for q if

$$v^T \nabla q(x) < 0.$$

Looking at Taylor expansion around x :

$$q(x+tv) = q(x) + \underbrace{t \nabla q(x)^T v}_{< 0 \text{ when } t > 0} + o(t)$$

and $\nabla q(x)^T v < 0$.

$\Rightarrow$ means there is a step length $t > 0$ for which $q(x+tv) < q(x)$.

Now look at previous step $\hat{x}_k$, which must solve:

$$\min_{\hat{x} \in \text{span}\{p_0, \dots, p_{k-1}\}} \frac{1}{2} \hat{x}^T A \hat{x} - \hat{x}^T b \Leftrightarrow$$

$$v^T (b - Ax_k) = 0, \\ \forall v \in \text{span}\{p_0, \dots, p_{k-1}\}$$

$(\Rightarrow)$

$$p_j^T (b - Ax_k) = 0 \\ \text{for } j = 0, \dots, k-1.$$

We also have:

$$x_{k+1} = x_k + \alpha_k p_k$$

$$0 = (Ax_{k+1} - b)^T p_j = (Ax_k - b)^T p_j + \alpha_k p_k^T A p_j \\ \text{for } j = 0, \dots, k-1$$

$\Rightarrow$ p_k is A -orthogonal to previous k directions

A -orthogonal means orthogonality w.r.t. inner product:

$$\langle u, v \rangle_A = \langle u, Av \rangle = u^T A v.$$

It is not hard to check $\langle u, v \rangle_A$ is indeed an inner product when A is s.p.d.

So we can obtain p_k be A -orthogonalizing family of vectors $\{r_0, r_1, \dots, r_{k-1}, r_k\}$.

Use Gram-Schmidt orthogonalization.

$$\begin{aligned} p_0 &= r_0 \\ p_k &= r_k - \sum_{j=0}^{k-1} \frac{r_k^T A p_j}{p_j^T A p_j} p_j \\ &= \perp \text{ proj w.r.t. } \langle \cdot, \cdot \rangle_A \text{ inner prod of } r_k \text{ along } p_j. \end{aligned}$$

Now take the optimal step size along this direction:

$$\alpha_k = \frac{p_k^T (b - A z_k)}{p_k^T A p_k} = \frac{p_k^T r_k}{p_k^T A p_k}$$

Gram Schmidt + this choice of step size is essentially a preliminary version of CG.

However in this version:

- cost of iteration grows linearly with iteration number.
- storage needed " " " " " .

Fortunately the Gram-Schmidt orthogonalization takes a simpler form if we look closely at subspaces we use.

The subspace that CG uses to look for new direction is: (54)

28

$$\begin{aligned}\text{span}\{p_0, p_1, \dots, p_{k-1}, r_k\} &= \text{span}\{p_0, \dots, p_{k-1}, p_k\} \\ &= \text{span}\{r_0, \dots, r_{k-1}, r_k\} \\ &= \mathcal{K}_{k+1}(A, r_0) \\ &= \underline{\text{Krylov subspace}}\end{aligned}$$

why?

$$r_0 = b - Ax_0 \in \mathcal{K}_1(A, r_0)$$

If $r_{k-1} = b - Ax_{k-1} \in \mathcal{K}_k(A, r_0)$ then

$$\begin{aligned}r_k &= b - Ax_k = b - A(x_{k-1} + \alpha p_{k-1}) \\ &= \underbrace{b - Ax_{k-1}}_{= r_{k-1} \in \mathcal{K}_k} - \underbrace{\alpha A p_{k-1}}_{\substack{\in \mathcal{K}_k \\ \in \mathcal{K}_{k+1}}}\end{aligned}$$

$$\Rightarrow \boxed{r_k \in \mathcal{K}_{k+1}}$$

Conjugate Gradient forms part of a large family of iterative solvers called Krylov Subspace methods.

Why did we introduce these subspaces?

Recall optimality conditions, now written with Krylov subspaces.

$$r_{k+1}^T v = 0, \quad \forall v \in \mathcal{K}_{k+1}(A, r_0)$$

$$r_k^T v = 0, \quad \forall v \in \mathcal{K}_k(A, r_0)$$

$$\text{Take } p_i \in \mathcal{K}_{k-1}(A, r_0) \Rightarrow A p_i \in \mathcal{K}_k(A, r_0)$$

$$\Rightarrow r_k^T A p_i = 0, \quad \text{for } i = 0, \dots, k-2$$

This orthogonality greatly simplifies Gram Schmidt and reduces sum to only one term:

$$p_k = r_k - \underbrace{\frac{r_k^T A p_{k-1}}{p_{k-1}^T A p_{k-1}}}_{\beta_k} p_{k-1}$$

- Storage and cost of iteration remains the same regardless of k .
- One needs roughly only 5 vectors of length n .
- One requires only knowledge of action of A on a vector.
 $\Rightarrow$ no need to know entries of A as in direct methods.
 A could be even specified via a "black-box" code.

In the classical formulation of CG α_k and β_k take different forms, which can be obtained by applying optimality conditions:

$$\alpha_k = \frac{\|r_k\|^2}{p_k^T A p_k}$$

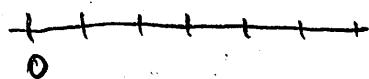
$$\beta_k = \frac{\|r_{k+1}\|^2}{\|r_k\|^2}$$

Convergence of CG:

- If A is s.p.d. CG converges in at most n steps, where n = dimension of A .
- In general convergence of CG depends on how many "clusters" of eigenvalues does A have.

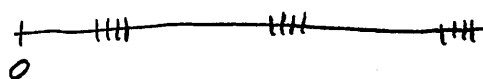
Spectrum of A (all eigenvalues)

(56)
30



slow convergence

(S1)



fast convergence in roughly
as many iterations as there are
clusters (here 3 iter)

(S2)

Preconditioning

To speed up convergence of CG (i.e. going from (S1) to (S2))
one solves system:

$$M^{-1}Ax = M^{-1}b, \text{ where } M \text{ is invertible}$$

$$\Leftrightarrow Ax = b$$

$M =$ (left) preconditioner

$=$ approximation of A that is easy to invert.

One can also think of right preconditioners.

$$A \underbrace{M^{-1}y}_=x = b \quad \Leftrightarrow Ax = b$$

For symmetric systems; we don't want to lose symmetry,
so letting

$$M = CC^T, M \text{ invertible:}$$

$$Ax = b \Leftrightarrow \underline{C^{-1}AC^{-T}} C^T x = C^{-1}b$$

s.p.d. if A s.p.d.

Note: there is no single way of choosing a preconditioner.
A preconditioner that works for a particular problem may not
work as well for another.

Here are some preconditioning techniques that work
(see Trefethen and Bau for a broader overview)

- incomplete LU or Cholesky factorization: drop elements that are below a threshold to get sparse LU (L) factors that give cheap systems to solve
- Jacobi: use diagonal or block-diagonal of A
- Discretization based: use solution to a coarse grid problem, a constant coeff problem, a periodic problem or solve problem in alternating directions (we will see this more closely when we get to ADI methods)

The final algorithm is as follows:

CONJUGATE GRADIENT (unpreconditioned)

$$x_0 = \text{given}$$

$$r_0 = b - Ax_0$$

$$p_0 = r_0$$

$$\text{for } n = 1, 2, 3, \dots$$

$$\alpha_n = \frac{r_{n-1}^T r_{n-1}}{p_{n-1}^T A p_{n-1}}$$

$$x_n = x_{n-1} + \alpha_n p_{n-1}$$

$$r_n = r_{n-1} - \alpha_n A p_{n-1} \quad \leftarrow \text{check convergence here by looking at } \|r_n\|$$

$$\beta_n = \frac{r_n^T r_n}{r_{n-1}^T r_{n-1}}$$

$$p_n = r_n + \beta_n p_{n-1}$$

and the preconditioned version is:

(58)

32

PRECONDITIONED CONJUGATE GRADIENT

$$x_0 = \text{given}$$

$$M = \text{preconditioner} = \text{given}$$

$$r_0 = b - Ax_0$$

$$\text{Solve } My_0 = r_0 \text{ for } y_0$$

$$p_0 = y_0$$

$$\text{for } n = 1, 2, 3, \dots$$

$$\alpha_n = \frac{r_{n-1}^T y_{n-1}}{p_{n-1}^T A p_{n-1}}$$

$$x_n = x_{n-1} + \alpha_n p_{n-1}$$

$$r_n = r_{n-1} - \alpha_n A p_{n-1} \quad \leftarrow \text{check convergence here.}$$

$$\text{Solve } My_n = r_n \text{ for } r_n. \quad \text{should be } y_n$$

$$\beta_n = \frac{y_n^T r_n}{y_{n-1}^T r_{n-1}}$$

$$p_n = y_n + \beta_n p_{n-1}$$

note: This is equivalent to applying CG to system

$$C^T A C^{-T} C^T x = C^{-T} b$$

where $M = C C^T = \text{preconditioner.}$

we will not see equivalence but notice:

$$y_n = M^{-1} r_n \Rightarrow r_n^T y_n = r_n^T M^{-1} r_n = (C^{-1} r_n)^T (C^{-1} r_n)$$

prec.
residual