



Topics in Numerical Analysis II Computational Inverse Problems

Lecturer: Bangti Jin (btjin@math.cuhk.edu.hk)

Chinese University of Hong Kong

November 5, 2022



Outline

1 Kaczmarz, randomized Kaczmarz

- Stochastic gradient descent

2 Krylov subspace methods



Review

decompose the system $Ax = y$ into

$$A_j x = y_j, \quad j = 1, \dots, \ell$$

- define an orthogonal projection $P_j : \mathbb{R}^n \rightarrow X_j = \{x : A_j x = y_j\}$

$$P_j x = x + A_j^\top (A_j A_j^\top)^{-1} (y_j - A_j x), \quad \forall x \in \mathbb{R}^n.$$

- define **sequential** proj. $P : \mathbb{R}^n \rightarrow \mathbb{R}^n$ by $P = P_\ell P_{\ell-1} \dots P_2 P_1$
- The Kaczmarz sequence $\{x_k\}_{k=0}^\infty$ is defined by Stephan Kaczmarz 1937

$$x_{k+1} = P x_k, \quad x_0 = 0$$

If $X = \cap_{j=1}^\ell X_j \neq \emptyset$, then the Kaczmarz sequence $\{x_k\}_{k=0}^\infty$ converges to the minimum norm solution $x^\dagger = A^\dagger y$ as $k \rightarrow \infty$.



Randomized Kaczmarz method

Stromer-Vershynin 2009: randomized Kaczmarz method

- choose the i th equation with probability prop. to $\|a_i\|^2$
- the method is a version of stochastic gradient descent

Let $x^\dagger$ be a solution to $Ax = y$. Then randomized Kaczmarz converges to x in expectation (with $\kappa(A) = \|A\|_F \|A^{-1}\|$):

$$\mathbb{E}[\|x_k - x^\dagger\|^2] \leq (1 - \kappa(A)^{-2})^k \|x_0 - x^\dagger\|^2$$

The method **converges exponentially fast** to $x^\dagger$, and the convergence rate depends only on scaled condition number $\kappa(A)$



convergence analysis via SVD for exact data

■ error $e_k = x_k - x^*$

■ error recursion

$$e_{k+1} = \left(I - \frac{a_{i_k} a_{i_k}^t}{\|a_{i_k}\|^2} \right) e_k.$$

$I - \frac{a_{i_k} a_{i_k}^t}{\|a_{i_k}\|^2}$ is an orthogonal projection operator

■ $\sigma_{\min} \|e\| \leq \|Ae\| \leq \sigma_1 \|e\|$



$$\begin{aligned}\|e_{k+1}\|^2 &= \|e_k\|^2 - \frac{2}{\|a_{i_k}\|^2} \langle a_{i_k}, e_k \rangle \langle a_{i_k}, e_k \rangle + \langle a_{i_k}, e_k \rangle^2 \frac{\|a_{i_k}\|^2}{\|a_{i_k}\|^4} \\ &= \|e_k\|^2 - \frac{1}{\|a_{i_k}\|^2} \langle a_{i_k}, e_k \rangle \langle a_{i_k}, e_k \rangle\end{aligned}$$

Upon noting the identity $\sum_{i=1}^n a_i a_i^t = A^t A$, taking expectation

$$\mathbb{E}[\|e_{k+1}\|^2 | e_k] \leq \|e_k\|^2 - \frac{1}{\|A\|_F^2} \langle e_k, A^t A e_k \rangle = \|e_k\|^2 - \frac{\|A e_k\|^2}{\|A\|_F^2}$$

$$+ \|A e\| \geq \sigma_{\min} \|e\| \Rightarrow$$

$$\mathbb{E}[\|e_{k+1}\|^2 | e_k] \leq \|e_k\|^2 - \frac{\sigma_{\min}^2 \|e_k\|^2}{\|A\|_F^2} = (1 - \kappa(A)^{-2}) \|e_k\|^2$$



Jiao-Jin-Lu, 2017, Inverse Problems

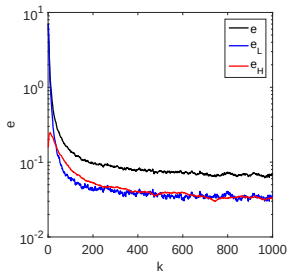
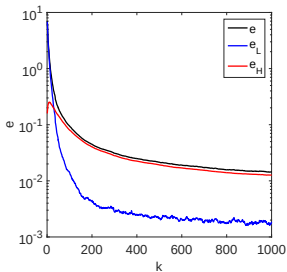
Let $c_1 = \frac{\sigma_L^2}{\|A\|_F^2}$ and $c_2 = \frac{\sum_{i=L+1}^r \sigma_i^2}{\|A\|_F^2}$. Then there hold

$$\mathbb{E}[\|P_L \mathbf{e}_{k+1}\|^2 | \mathbf{e}_k] \leq (1 - c_1) \|P_L \mathbf{e}_k\|^2 + c_2 \|P_H \mathbf{e}_k\|^2,$$

$$\mathbb{E}[\|P_H \mathbf{e}_{k+1}\|^2 | \mathbf{e}_k] \leq c_2 \|P_L \mathbf{e}_k\|^2 + (1 + c_2) \|P_H \mathbf{e}_k\|^2.$$

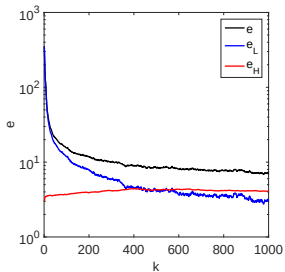
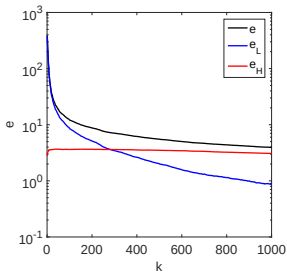


Example: phillips





Example: gravity





randomized Kaczmarz as SGD

solving the problem by approximately minimizing

$$J(x) = (2n)^{-1} \|Ax - y\|^2 = \frac{1}{n} \sum_{j=1}^n f_j(x), \quad f_j(x) = \frac{1}{2} ((a_j, x) - y_j)^2.$$

many different approaches:

- gradient descent (a.k.a. Landweber iteration)

$$x_{k+1} = x_k - \eta_k n^{-1} A^t (Ax_k - y)$$

acceleration by step-size, Nesterov trick, momentum, Anderson,

...

- **stochastic** gradient descent Robbins-Monro 1951 $f_j = \frac{1}{2} ((a_{i_k}, x) - b_{i_k})^2$

$$x_{k+1} = x_k - \eta_k ((a_{i_k}, x_k) - b_{i_k}) a_{i_k}$$

often: index i_k uniformly over the set $\{1, \dots, n\}$



RKM and SGD: the linear equation $Ax = y$ in least squares form:

$$\min_{x \in \mathbb{R}^n} \frac{1}{2n} \sum_{i=1}^n |(a_i, x) - y_i|^2 \triangleq f(x) = \sum_{i=1}^n \frac{\|a_i\|^2}{\|A\|_F^2} f_i(x),$$

with

$$f_i(x) = \frac{\|A\|_F^2}{2n\|a_i\|^2} |(a_i, x) - y_i|^2$$

and sampling probability of drawing i th row being Strohmaier, Vershynin 2009

$$p_i = \frac{\|a_i\|^2}{\|A\|_F^2}.$$

this is a valid probability distribution

$$\sum_{i=1}^n p_i = \sum_{i=1}^n \frac{\|a_i\|^2}{\|A\|_F^2} = 1$$



Needell et al 2016; Jiao-Jin-Lu 2017 Inverse Problems

RKM is a (weighted) SGD with a **constant stepsize** $\beta = n/\|A\|_F^2$.

$$\begin{aligned}x_{k+1} &= x_k - \beta \partial_x f_{i_k}(x_k) \\&= x_k - \frac{n}{\|A\|_F^2} \cdot \frac{\|A\|_F^2}{n\|a_{i_k}\|^2} ((a_{i_k}, x) - y_{i_k}) a_{i_k} \\&= x_k - ((a_{i_k}, x) - y_{i_k}) \frac{a_{i_k}}{\|a_{i_k}\|^2}\end{aligned}$$



The SGD has a long history in statistics

- Robbins-Monro 1950s, a few monographs Kushner-Yin 2003
various asymptotic distributional results ...
- revived interest in machine learning and signal processing
- accelerated variants represent the state of art solvers for large-scale problems: machine learning, inverse problems etc.



Asymptotic convergence

Question: Is SGD consistent (as $\delta \rightarrow 0$) ?

This is one of the basic requirements of a reconstruction scheme ...
since x_k^δ is random, which norm do we use ?

- strong convergence (mean squared error)

$$\lim_{\delta \rightarrow 0} \mathbb{E}[\|x_{k(\delta)}^\delta - x^\dagger\|^2] = 0 \quad ???$$

- weak convergence (easy!)

$$\lim_{\delta \rightarrow 0} \|\mathbb{E}[x_{k(\delta)}^\delta] - x^\dagger\|^2 = 0 \quad ???$$

- high-probability / almost sure ...



weak convergence

$$x_k = x_{k-1} - \eta_k ((a_{i_k}, x_{k-1}) - y_{i_k}) a_{i_k}$$

$\Rightarrow$

$$\begin{aligned}\mathbb{E}_k[x_k] &= x_{k-1} - \eta_k \mathbb{E}_k[((a_{i_k}, x_{k-1}) - y_{i_k}) a_{i_k}] \\ &= x_{k-1} - \eta_k n^{-1} \sum_{i=1}^n ((a_i, x_{k-1}) - y_i) a_i \\ &= x_{k-1} - \eta_k n^{-1} A^\top (Ax_{k-1} - y)\end{aligned}$$

taking full expectation

$$\mathbb{E}[x_k] = \mathbb{E}[x_{k-1}] - \eta_k n^{-1} A^\top (A\mathbb{E}[x_{k-1}] - y)$$

$\mathbb{E}[x_k]$ satisfies the Landweber iteration



Question:

$$\lim_{\delta \rightarrow 0} \mathbb{E}[\|x_{k(\delta)}^\delta - x^\dagger\|^2] = 0 \quad ???$$

Answer: **Yes** for the step size choice:

$$\eta_j = c_0 j^{-\alpha}, \quad \alpha \in (0, 1)$$

if

$$\lim_{\delta \rightarrow 0} k(\delta) = \infty \quad \text{and} \quad \lim_{\delta \rightarrow 0} k(\delta)^{\frac{1-\alpha}{2}} \delta = 0.$$



basic strategy of proof: by bias-variance decomposition

$$\begin{aligned}\mathbb{E}[\|x_k^\delta - x^\dagger\|^2] &= \underbrace{\|\mathbb{E}[x_k^\delta] - x^\dagger\|^2}_{\text{mean}} + \underbrace{\mathbb{E}[\|\mathbb{E}[x_k^\delta] - x_k^\delta\|^2]}_{\text{variance}} \\ &\leq 2 \underbrace{\|\mathbb{E}[x_k^\delta - x_k]\|^2}_{\text{propagation error}} + 2 \underbrace{\|\mathbb{E}[x_k] - x^\dagger\|^2}_{\text{approximation error}} + \underbrace{\mathbb{E}[\|\mathbb{E}[x_k^\delta] - x_k^\delta\|^2]}_{\text{stochastic error}}\end{aligned}$$

- propagation error and approximation error are well understood (i.e., Landweber method)
- stochastic error needs to be controlled ...



how to control the stochastic error

$$\mathbb{E}[\|x_{k+1}^\delta - \mathbb{E}[x_{k+1}^\delta]\|^2] \leq \sum_{j=1}^k \eta_j^2 \|B^{\frac{1}{2}} \Pi_{j+1}^k(B)\|^2 \mathbb{E}[\|Ax_j^\delta - y^\delta\|^2]$$

$\Rightarrow$ bound the expected residual properly: for any $\ell \geq 2$

$$\mathbb{E}[\|Ax_{k+1}^\delta - y^\delta\|^2] \leq ck^{-\min(\ell\alpha, 2(1-\alpha), (2p+1)(1-\alpha))} \ln^\ell k + c'\delta^2 \ln^2 k.$$

with the parameter p (source condition)

$$B^p w = x^\dagger - x_1$$

regularity condition on the initial error: the larger is p , the smoother is the solution ...



limiting cases:

- no solution regularity $p = 0$

$$\mathbb{E}[\|Ax_{k+1}^\delta - b^\delta\|^2] \leq ck^{-\min(\ell\alpha, 1-\alpha)} \ln^\ell k + c'\delta^2 \ln^2 k.$$

- very good regularity

$$\mathbb{E}[\|Ax_{k+1}^\delta - b^\delta\|^2] \leq ck^{-\min(\ell\alpha, 2(1-\alpha))} \ln^\ell k + c'\delta^2 \ln^2 k.$$

$\Rightarrow$ randomness **restricts the best possible decay** of the residual error estimates

$$\mathbb{E}[\|x_k^\delta - x^\dagger\|^2] \leq ck^{-\min(\ell\alpha, 1-\alpha, 2p(1-\alpha))} \ln^\ell k + c\delta^2 k^{1-\alpha}.$$

by choosing proper $k(\delta) \Rightarrow$ error estimates

SGD is consistent with convergence rate (if we know how to stop)!

challenge: Nobody knows how to stop !



what is beyond ?

- SGD / randomized Kaczmarz is regularizing in the mean squares error sense, but optimality not well understood
- stopping criterion is not well understood
- convergence in other modes (a.s., high-probability) not well understood



Conjugate gradient method

large linear systems of the form $Ax = y$, $A \in \mathbb{R}^{n \times n}$

- Krylov subspace methods are methods of choice
 - well known version: the conjugate gradient method
 - generalized minimal residual method (GMRES)
 - biconjugate gradient method (BiCG) (stabilized)
 - ...
- to approx. x by a linear combination of vectors $u, Au, A^2u, \dots$
- If Av is cheap (e.g., A is sparse), Krylov subspace methods are especially efficient.

goal: introduce the basic idea behind CG and show its application



assumptions on A and related inner-product

- the matrix $A \in \mathbb{R}^{n \times n}$ is s.p.d.

$$A^\top = A \quad \text{and} \quad u^\top A u > 0, \quad \forall u \neq 0$$

(A is invertible, the inverse $A^{-1} \in \mathbb{R}^{n \times n}$ is also s.p.d.)

- A -dependent inner product and the corresponding norm:

$$(u, v)_A = u^\top A v \quad \text{and} \quad \|u\|_A = (u, u)_A^{\frac{1}{2}}$$

$\Rightarrow (\cdot, \cdot)_A : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ is an inner product, and $\|\cdot\|_A$ is a norm.

[This setting generally does apply to inverse problems directly!]



- $x^* = A^{-1}y \in \mathbb{R}^n$ be the unique solution of $Ax = y$
- the error and the residual for some approximation x by

$$e = x^* - x \quad \text{and} \quad r = y - Ax = Ae$$

- $\phi : \mathbb{R}^n \rightarrow \mathbb{R}$ by

$$\phi(x) = \|e\|_A^2 = (e, Ae) = (r, A^{-1}r) = \|r\|_{A^{-1}}^2$$

- $\|\cdot\|_A$ is a norm $\Rightarrow \phi(x) \geq 0$ and $= 0$ iff

$$e = 0 \quad \Leftrightarrow \quad x = x^*$$

- observation: minimizing ϕ is equivalent to solving $Ax = y$



minimizing ϕ in a given direction d

- evaluating ϕ requires x^* , or $A^{-1} \Rightarrow$ **unrealistic**
- fix $x_0 \in \mathbb{R}^n$ and direction $0 \neq d_0 \in \mathbb{R}^n$, minimizing ϕ over the line

$$D = \{x = x_0 + \alpha d_0 : \alpha \in \mathbb{R}\}$$

does not require x^* or A^{-1}



The function $\alpha \mapsto \phi(x_0 + \alpha d_0)$, $\mathbb{R} \rightarrow \mathbb{R}$ attains its minimum at

$$\alpha = \alpha_0 := \frac{(d_0, r_0)}{\|d_0\|_A^2} = \frac{(d_0, r_0)}{(d_0, Ad_0)}$$

with $r_0 = y - Ax_0$.

Key: the computation does not require A^{-1} or x^* !



The residual r corresponding to $x = x_0 + \alpha d_0$ is

$$r = y - Ax = y - Ax_0 - \alpha Ad_0 = r_0 - Ad_0$$

and hence

$$\begin{aligned}\phi(x) &= (r, A^{-1}r) \\ &= (r_0 - \alpha Ad_0, A^{-1}(r_0 - \alpha Ad_0)) \\ &= \alpha^2(d_0, Ad_0) - 2\alpha(d_0, r_0) + (r_0, A^{-1}r_0),\end{aligned}$$

quadratic function in α

$\Rightarrow$ the minimum is at the unique zero of the derivative in α , i.e., $\alpha = \alpha_0$.



sequential minimization

- Given a sequence of nonzero search directions $(d_k)_k$
 $\Rightarrow$ a sequence of approx. by iteratively minimizing ϕ on
 $\{x_k + \alpha d_k : \alpha \in \mathbb{R}\}$

$$x_{k+1} = x_k + \alpha_k d_k, \quad \alpha_k = \frac{(d_k, r_k)}{(d_k, A d_k)}, \quad k = 0, 1, \dots$$

r_k is the residual for the k th iterate, i.e.,

$$r_k = y - A x_k$$

- $(\phi(x_k))$ is a decreasing sequence: $\phi(x_{k+1}) \leq \phi(x_k)$
- the choice of the search directions (d_k) is subtle !



- one first idea is to choose Cauchy 1847

$$d_k = -\nabla\phi(x_k) = 2(y - Ax_k) = 2r_k$$

gives the direction of the steepest descent (Landweber):
ineffective

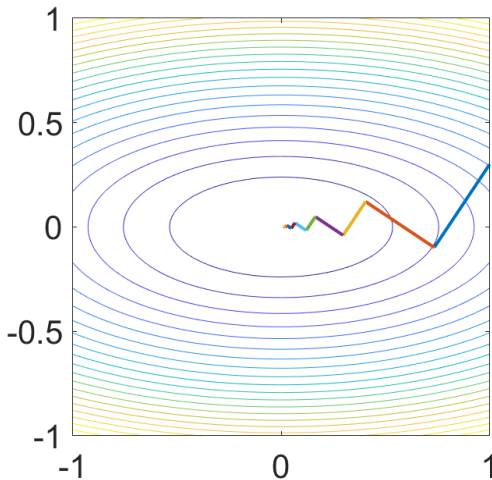
- often not giving a sequence (x_k) fast convergent to $x^* = A^{-1}y$

Example

$$A = \begin{bmatrix} 1 & 0 \\ 0 & 5 \end{bmatrix}, \quad y = \begin{bmatrix} 0 \\ 0 \end{bmatrix}, \quad \phi(x) = (x^{(1)})^2 + 5(x^{(2)})^2$$



convergence plot of steepest descent





minimizing ϕ over a hyperplane

- $(d_0, \dots, d_k)$ a set of linearly independent search directions
- finding the minimizer of ϕ on the hyperplane

$$\mathcal{D}_k = \{x = x_0 + D_k h : h \in \mathbb{R}^{k+1}\}$$

with $x_0 \in \mathbb{R}^n$ and $D_k = [d_0 \ d_1 \ \dots \ d_k] \in \mathbb{R}^{n \times (k+1)}$

- the function $h \mapsto \phi(x_0 + D_k h)$, $\mathbb{R}^{k+1} \rightarrow \mathbb{R}$ attains its minimum at

$$h^* = (D_k^\top A D_k)^{-1} D_k^\top r_0$$



- $D_k^\top AD_k \in \mathbb{R}^{(k+1) \times (k+1)}$ is invertible:

$$D_k^\top AD_k z = 0 \quad \Rightarrow \quad z^\top D_k^\top AD_k z = 0 \quad \Rightarrow \quad D_k z = 0$$

the columns of D_k are linearly independent $\Rightarrow z = 0$

- $\ker(D_k^\top AD_k) = \{0\}$, and $(D_k^\top AD_k)^{-1}$ exists.
- The residual r for $x = x_0 + D_k h$ satisfies

$$r = y - A(x_0 + D_k h) = r_0 - AD_k h$$

and thus

$$\begin{aligned} \phi(x_0 + D_k h) &= (r_0 - AD_k h)^\top A^{-1} (r_0 - AD_k h) \\ &= h^\top D_k^\top AD_k h - 2(r_0, D_k h) + (r_0, A^{-1} r_0) \end{aligned}$$



- the matrix $D_K^\top A D_K$ of the quadratic form in h is positive definite
- the unique zero of the gradient of $\phi(x_0 + D_K h)$ in h :

$$h_* = (D_K^\top A D_K)^{-1} D_K^\top r_0$$

is the unique minimizer of $\phi(x_0 + D_K h)$ over h



A-conjugate search direction

comparison between seq. min v.s. subspace min.

- minimizing ϕ over $\mathcal{D}_k$ involves inverting a $D_k^\top A D_k$ matrix
 $\Rightarrow$ **computationally not attractive**
- minimizing ϕ seq. in $(d_j)_{j=0}^k$ often does not coincide with minimizing over $\mathcal{D}_k$
- seq. min. reproduces the minimizer over $\mathcal{D}_k$ if $(d_j)_{j=0}^k$ chosen cleverly $\Leftarrow$ **A-conjugacy**



A -conjugacy

- the nonzero vectors $(d_0, \dots, d_k)$ are A -conjugate if

$$(d_i, d_j)_A = 0, \quad i \neq j$$

A -conjugate if **orthogonal w.r.t. the inner product $(\cdot, \cdot)_A$**

- The A -conjugacy condition for the matrix $D_k = [d_0 \ \dots \ d_k]$:

$$D_k^\top A D_k = \text{diag}(s_0, s_1, \dots, s_k) \in \mathbb{R}^{(k+1) \times (k+1)}$$

with $s_j = (d_j, d_j)_A > 0$ ($\because A$ is s.p.d.)



implications of A -conjugate directions

Given x_0 , let $(d_j)_{j=0}^k \subset \mathbb{R}^n$ are nonzero and A -conjugate. The seq. minimizers of ϕ

$$x_{j+1} = x_j + \alpha_j d_j, \quad \text{with } \alpha_j = \frac{(d_j, r_j)}{(d_j, d_j)_A}, \quad j = 0, 1, \dots, k,$$

is a minimizer of ϕ on the hyperplane $\mathcal{D}_k$:

$$x_{k+1} = x_0 + D_k h_* = x_0 + D_k (D_k^\top A D_k)^{-1} D_k r_0.$$



- Let $a_j = (\alpha_0, \dots, \alpha_j)^\top \in \mathbb{R}^{j+1}$. Then we have

$$x_j = x_0 + \sum_{i=0}^{j-1} \alpha_i d_i = x_0 + D_{j-1} a_{j-1}, \quad j = 1, \dots, k+1$$

- the corresponding residual r_j is

$$r_j = y - Ax_j = (y - Ax_0) - AD_{j-1} a_{j-1} = r_0 - AD_{j-1} a_{j-1}$$

- d_j is A -conjugate to $(d_0, \dots, d_{j-1}) \Rightarrow$

$$(d_j, r_j) = (d_j, r_0) - \underbrace{(d_j, AD_{j-1} a_{j-1})}_{=0} = (d_j, r_0)$$

- α_j depends only on r_0 and d_j

$$\alpha_j = \frac{(d_j, r_j)}{(d_j, d_j)_A} = \frac{(d_j, r_0)}{(d_j, d_j)_A}, \quad j = 0, \dots, k$$



- $(d_0, \dots, d_k)$ are A -conjugate $\Rightarrow$

$$\begin{aligned}(D_k^\top A D_k)^{-1} &= (\text{diag}(d_0^\top A d_0, \dots, d_k^\top A d_k))^{-1} \\ &= \text{diag}((d_0^\top A d_0)^{-1}, \dots, (d_k^\top A d_k)^{-1}) \in \mathbb{R}^{(k+1) \times (k+1)}\end{aligned}$$

- equivalence relation

$$h_* = (D_k A D_k)^{-1} D_k^\top r_0 = (D_k^\top A D_k)^{-1} \begin{bmatrix} (d_0, r_0) \\ \vdots \\ (d_k, r_0) \end{bmatrix} = \begin{bmatrix} \alpha_0 \\ \vdots \\ \alpha_k \end{bmatrix} = a_k$$

$$\Rightarrow a_k = h_*,$$

$$x_{k+1} = x_0 + D_k a_k = x_0 + D_k h_*$$

seq. minimizer = minimizer over $\mathcal{D}_k$



A-conjugate directions $\Rightarrow$ orthogonality w.r.t. residuals

If the nonzero search directions $(d_j)_{j=0}^k \subset \mathbb{R}^n$ are A -conjugate, then $r_{k+1} = y - Ax_{k+1}$ satisfies

$$r_{k+1} \perp \text{span}(d_0, \dots, d_k)$$

■ $x_{k+1} = x_0 + D_k h_* \Rightarrow$

$$r_{k+1} = (y - Ax_0) - AD_k h_* = r_0 - AD_k h_*$$

■ $h_*^\top = ((D_k^\top AD_k)^{-1} D_k^\top r_0)^\top = r_0^\top D_k (D_k^\top AD_k)^{-1} \Rightarrow$

$$\begin{aligned} [(r_{k+1}, d_0), \dots, (r_{k+1}, d_k)] &= r_{k+1}^\top D_k \\ &= r_0^\top D_k - h_*^\top D_k^\top AD_k = 0 \end{aligned}$$



Construct A -conjugate search directions?

CG constructs A -conjugate search directions using Krylov subspaces

Definition

The k th Krylov subspace of A with $r_0 = y - Ax_0$ is defined by

$$\mathcal{K}_k = \mathcal{K}_k(A, r_0) = \text{span}(\{r_0, \dots, A^{k-1}r_0\}), \quad k = 1, 2, \dots,$$

Note that $A(\mathcal{K}_k) \subset \mathcal{K}_{k+1}$, and $\mathcal{K}_{k-1} \subset \mathcal{K}_k$.



construct A -conjugate directions inductively

given a set of A -conjugate vectors, either find a new A -conjugate direction or the current iterate is the global minimizer.

- 1 Choose an initial guess $x_0 \in \mathbb{R}^n$
- 2 If $r_0 = y - Ax_0 = 0$, then $x^* = x_0$. Otherwise set $d_0 = r_0$
- 3 given nonzero A -conjugate search directions $(d_j)_{j=0}^{k-1}$ such that

$$\mathcal{K}_m = \text{span}(\{d_j\}_{j=0}^{m-1}) = \text{span}(\{r_j\}_{j=0}^{m-1}), \quad m = 1, \dots, k$$

with $r_j = y - Ax_j$: the residual in the seq. min

- $r_k \equiv 0$: converged
- $r_k \neq 0$: find A -conjugate direction $d_k \neq 0$ (and the relation remains valid for $k+1$)



- assuming $r_k \neq 0$ + the identity

$$r_k = y - Ax_k = y - A(x_{k-1} + \alpha_{k-1}d_{k-1}) = r_{k-1} - \alpha_{k-1}Ad_{k-1}$$

and $r_{k-1}, d_{k-1} \in \mathcal{K}_k \Rightarrow$ the residual r_k belongs to $\mathcal{K}_{k+1}$

- r_k is orthogonal to $(d_j)_{j=0}^{k-1}$, which span $\mathcal{K}_k$ and belong to $\mathcal{K}_{k+1} \Rightarrow$

$$\mathcal{K}_{k+1} = \text{span}(d_0, \dots, d_{k-1}, r_k) = \text{span}(r_0, \dots, r_{k-1}, r_k)$$

- find the new direction d_k in the form

$$d_k = r_k + \beta_{k-1}d_{k-1}, \quad \beta_{k-1} \in \mathbb{R}.$$

$\Rightarrow d_k$ belongs to $\mathcal{K}_{k+1}$ and

$$\mathcal{K}_{k+1} = \text{span}(d_0, \dots, d_{k-1}, r_k) = \text{span}(d_0, \dots, d_{k-1}, d_k)$$

A-conjugacy condition ??



- choose β_{k-1} so that for $j = 0, \dots, k-1$

$$\begin{aligned} d_j^\top Ad_k &= d_j^\top Ar_k + \beta_{k-1} d_j^\top Ad_{k-1} \\ &= (Ad_j)^\top r_k + \beta_{k-1} d_j^\top Ad_{k-1} = 0??? \end{aligned}$$

- $\{d_0, \dots, d_{k-2}\} \subset \mathcal{K}_{k-1} \Rightarrow$

$$\{Ad_0, \dots, Ad_{k-2}\} \subset \mathcal{K}_k = \text{span}(d_0, \dots, d_{k-1})$$

$\Rightarrow$ the vectors $\{Ad_0, \dots, Ad_{k-2}\}$ are orthogonal to d_k .

- need only to verify $j = k-1$ for A -conjugacy
- Solving the equation for β_{k-1} yields

$$d_k = r_k + \beta_{k-1} d_{k-1}, \quad \beta_{k-1} = -\frac{(d_{k-1}, Ar_k)}{(d_{k-1}, d_{k-1})_A}$$



complete algorithm.

- 1: Choose x_0
- 2: Set $k = 0$, $r_0 = y - Ax_0$, $d_0 = r_0$;
- 3: **while** the stopping rule is not satisfied **do**
- 4: $\alpha = \frac{(d_k, r_k)}{(d_k, d_k)_A}$;
- 5: $x_k = x_{k-1} + \alpha d_k$;
- 6: $r_k = r_{k-1} - \alpha A d_k$;
- 7: $\beta = -\frac{(d_k, r_k)_A}{(d_k, d_k)_A}$;
- 8: $d_{k+1} = r_k + \beta d_k$;
- 9: $k = k + 1$;
- 10: **end while**



- In practice, the algorithm is presented in a slightly different form.
- If not yet converged at x_k , use the following formula: $r_k \perp d_{k-1} \Rightarrow$

$$(d_k, r_k) = (r_k + \beta_{k-1} d_{k-1}, r_k) = \|r_k\|^2$$

resulting in

$$\alpha_k = \frac{\|r_k\|^2}{(d_k, d_k)_A}$$

- Since $r_{k+1} \perp \text{span}(d_0, \dots, d_k) = \mathcal{K}_{k+1} \ni r_k$, we have

$$\|r_{k+1}\|^2 = (r_{k+1}, r_k - \alpha_k A d_k) = -\frac{\|r_k\|^2}{(d_k, d_k)_A} (r_{k+1}, d_k)_A = \beta_k \|r_k\|^2$$

Solving for β_k

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

the formulae for α_k and $\beta_k \Rightarrow$ the standard form of the method.



- 1: Choose x_0
- 2: Set $k = 0$, $r_0 = y - Ax_0$, $d_0 = r_0$;
- 3: **while** the stopping rule is not satisfied **do**
- 4: $\alpha = \frac{(r_k, r_k)}{(d_k, d_k)_A}$;
- 5: $x_{k+1} = x + \alpha d_k$;
- 6: $r_{k+1} = r_k - \alpha A d_k$;
- 7: $\beta = \frac{(r_{k+1}, r_{k+1})}{(r_k, r_k)}$;
- 8: $d_{k+1} = r_{k+1} + \beta d_k$;
- 9: $k = k + 1$;
- 10: **end while**



CG for linear inverse problems

$$Ax = y$$

with $A \in \mathbb{R}^{n \times n}$: s.p.d.

- CG obtains the exact solution after at most n iterations
up to rounding errors only, can be delicate in low-prec.
- in practice, CG typically converges much quicker
- for ill-posed problems, terminate the iteration very carefully, in order to avoid overfitting to the noise
- extremely cautious, because CG often converges very fast



general ill-posed problem

$$Ax = y$$

with $A \in \mathbb{R}^{m \times n}$, $y \in \mathbb{R}^m$

- $m = n$ and A might be s.p.d: apply CG directly.
- generally, consider the normal equation

$$A^\top Ax = A^\top y$$

solving the problem in the least-squares sense.

- the matrix $A^\top A$ is symmetric, p.s.d.
- the conditions of CG are **almost satisfied**
- stopping with discrepancy principle: terminate when

$$\|y - Ax_k\| \leq \delta$$



backward heat problem:

- The matrix $A = e^{TB}$ is symmetric, positive semidefinite
- appears reasonable to apply CG directly
- add a small amount of noise
- discrepancy principle: terminate the iteration when

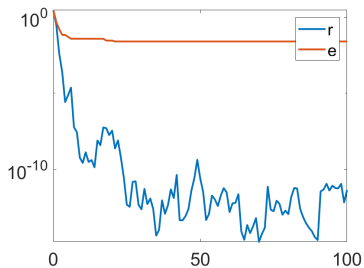
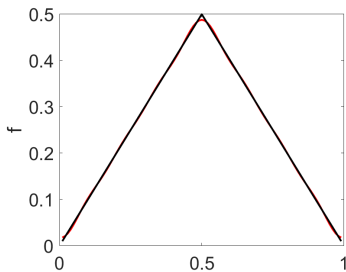
$$\|y - Ax_k\| \leq \delta$$

alternatively, apply CG to the normal equation

- the use of the normal equation makes the algorithm more stable and the sol. looks nicer for noisy data! (However, it is slower!)

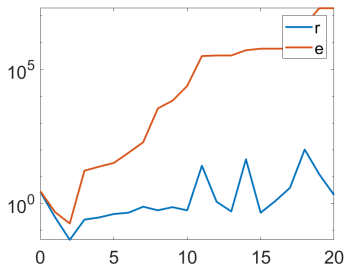
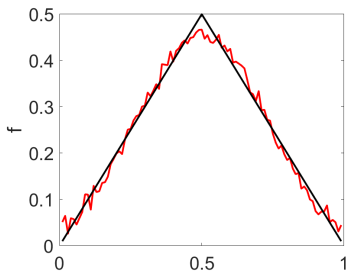


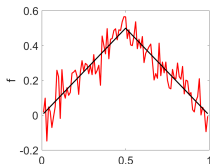
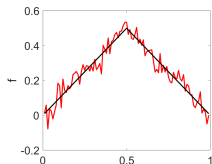
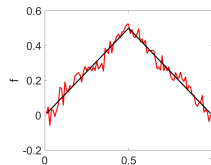
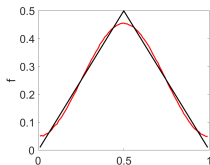
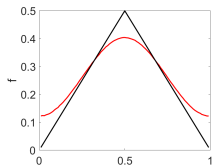
result for exact data, direct





result for noisy data, direct v.s. normal







result for noisy data, normal form, $x = A^T z$

