

# Lecture #12



→ We've been thru the barriers, culminating in a derivation of KKT. mentioned brackets  
another deriv based on  
in  $f(x) + \max_i f_i(x)$

Example:  $\min_x \sum_i \log(1 + e^{a_i^T x})$  s.t.  $x^T b = 1$

NNLS:  $\min \frac{1}{2} \|Ax - b\|^2$  s.t.  $x \geq 0$

$L(x, \lambda) = \frac{1}{2} \|Ax - b\|^2 - \lambda^T x$

$\phi(x) = \max \{f(x) - f^*, h_1(x), \dots, h_m(x)\}$

Finding direct  $\left[ \begin{array}{l} \frac{\partial L}{\partial x} \cdot A^T(Ax - b) - \lambda = 0 \\ \lambda^T x = 0 \\ x \geq 0, \lambda \geq 0 \end{array} \right. \quad \begin{array}{l} \text{--- (1)} \\ \text{(nonlinear!)} \end{array}$   
solutions not easy?

→ Usually, KKT equations not easy to solve directly, iterative approaches are taken

→ Even more simple: consider a function  $f(x)$ :

$f(x) = \frac{1}{2} x^T A x + b^T x$ ,  $A \succeq 0$

$\nabla f(x) = Ax + b = 0 \Rightarrow$  solve linear system  $Ax = -b$ .

• If  $b \in \text{Ran}(A)$ , system has a solution.

Direct methods: Chol. decomp of  $A$ :  $\frac{1}{3} n^3$  flops

Iterative methods: (Vast area within numerical linear algebra since well over 40 years!)

Many key opt. algorithmic ideas originally stem from trying to solve  $Ax = b$ !

→ At an abstract level:

$\min f(x)$

(1)  $\Rightarrow \nabla f(x) = 0 \Rightarrow x = (\nabla f)^{-1}(0) = \nabla f^*$

(2)  $\nabla f(x) = 0 \Rightarrow -\alpha \nabla f(x) = 0$  for  $\alpha > 0$  too  
 $\Rightarrow x - \alpha \nabla f(x) = x$  must hold

Suggests an algorithm!  $x_{k+1} = x_k - \alpha_k \nabla f(x_k)$ !

The file that someone asked me to print  
 really looked like a binary file.  
 So I printed this page instead.  
 Since printing a binary file is usually a waste  
 of paper and other resources, the spooling  
 software refuses to honor your request.  
 For comments or questions, send email to Help@cs.cmu.edu  
 Or visit <http://www.cs.cmu.edu/~help>.

1.5

$$-\alpha \nabla f(x) = 0$$

$$\Rightarrow -\alpha \nabla f(x) = 0 \quad \text{show that if } S \text{ is invariant}$$

we can find  $x - \alpha (\nabla^2 f(x))^{-1} \nabla f(x)$ . Newton method.

# Other "creative" possibilities sometimes

$$\begin{aligned} f(x) - g(x) \\ \nabla f(x) - \nabla g(x) = 0 \\ \nabla f(x) = \nabla g(x) \end{aligned}$$

Numerous  
 possibilities to  
 play with them

(soon we'll see a practical choice)

→ this is the "algebraic" derivation of Gradient descent

⑤

→ Say  $f(x)$  is non-diff. Then,

$$0 \in \partial f(x)$$

Now, actually,  $(\partial f)^{-1}(0)$  is a multivalued map: a set.

Now about

$$x \in x + \alpha \partial f(x)$$

as before? less clear. because this generates!

$$\underline{x_{k+1} \in x_k + \alpha_k \partial f(x_k)}$$

(With careful choice of  $\alpha_k$ , using any  $g_k \in \partial f(x_k)$  works!)

$$\rightarrow 0 \in \partial f(x) \Rightarrow x \in x + \alpha \partial f(x)$$

$$\Rightarrow (Id + \alpha \partial f)(x) \ni x$$

$$\Rightarrow \underline{x = (Id + \alpha \partial f)^{-1}(x)}$$

$$\Rightarrow \underline{x_{k+1} = (Id + \alpha \partial f)^{-1}(x_k)}$$

(Fixed-point iterat.)

This is actually Rockafellar's famous "proximal point algorithm" → A method that ultimately includes all sorts of popular iterations (as a proper choice of  $\partial f$ )

→ Q: A funny no reason for above "guess" to work → but this as a good starting point.

Q: Say we have  $\phi(x) = f(x) + g(x)$

(3)

(a)  $x_{k+1} = x_k - \alpha_k \nabla \phi(x_k)$

Seems to hold if  $\nabla \phi(x_k) = \nabla f(x_k) + \nabla g(x_k)$

(b) Suppose  $\phi(x) = f(x) - g(x)$  where both  $f, g$  are

How should we go about minimizing such a  $\phi$ ?



Key idea is: • Start with original frame and a guess, say  $x_k$ :  
Opt • Build an approx. to  $\phi(x)$ , centered at  $x_k$ , with the aim that approx is easy to optimize; call it  $m(x; x_k)$   
 • Obtain  $x_{k+1} = \arg \min m(x; x_k)$

• Gradient descent:  $\phi(x+\delta) \approx \underbrace{\phi(x) + \langle \nabla \phi(x), \delta \rangle}_{m(x, \delta)} + \frac{1}{2\alpha} \|\delta\|^2$

$x_{k+1} = \arg \min_{\delta} x_k + \delta$

$\delta = \arg \min_{\delta} \phi(x) + \langle \nabla \phi(x), \delta \rangle + \frac{1}{2\alpha} \|\delta\|^2$   
 $= \nabla \phi(x) + \frac{1}{\alpha} \delta = 0$   
 $\Rightarrow \delta = -\alpha \nabla \phi(x)$

$\delta \approx \boxed{x_{k+1} = x_k - \alpha_k \nabla \phi(x_k)}$  : G.D.

•  $\phi(x) = f(x) - g(x)$ ;  $f, g$  are

$g(x+\delta) \geq g(x) + \langle \nabla g(x), \delta \rangle$

$\delta \approx \phi(x+\delta) = f(x+\delta) - g(x+\delta)$

$\leq \underbrace{f(x+\delta) - g(x) - \langle \nabla g(x), \delta \rangle}_{m(x, \delta)}$

$x_{k+1} = x_k + \delta$   
 $\delta = \arg \min_{\delta} f(x+\delta) - \langle \nabla g(x), \delta \rangle$

(3.5)

Example of CCLSuppose we wish to minimize

$$\Phi(X) = \min_{X > 0} \sum_i \log \frac{\det(X \frac{a_i}{2})}{\sqrt{\det(X) \det(a_i)}}$$

This problem is called Frédet-mean,

$$\begin{aligned} \Phi(X) &= \sum_i \left[ \log \left( X \frac{a_i}{2} \right) - \frac{1}{2} \log \det(a_i) \right] \\ &= \sum_i \left[ \underbrace{\frac{1}{2} \log \det(X a_i)}_{f(X)} - \underbrace{\left( -\log \left( X \frac{a_i}{2} \right) \right)}_{g(X)} \right] \end{aligned}$$

Obtain linearization:

$$m(X|X) = \sum_i -\frac{1}{2} \log \left( X \frac{a_i}{2} \right) - \frac{1}{2} \langle \left( X \frac{a_i}{2} \right)^{-1}, s \rangle$$

we get  $\sum_i \left( -\frac{1}{2} (X \frac{a_i}{2})^{-1} - \frac{1}{2} (X \frac{a_i}{2})^{-1} \right) = 0$

$$\Rightarrow n (X \frac{a_i}{2})^{-1} = \sum_i (X \frac{a_i}{2})^{-1}$$

$$\Rightarrow (X \frac{a_i}{2})^{-1} = \frac{1}{n} \sum_i (X \frac{a_i}{2})^{-1}$$

$$\Rightarrow X_{\text{new}} = X_{\text{old}} = \left[ \quad \right]^{-1}$$

# Does this work? Converse? Etc?

Actually this particular example always has global optimum!

~~# Beyond CC~~

~~# However, Not always do.~~

# The art of using CCP lies in  
decomposing a given non convex  
function into well parts that allow  
easier optimization.

# It ensures  $\downarrow$  etc.

# Reading assignment: "On the convergence of the Convex-Concave Procedure"  
Srikanth Sridharan, L'Landon

# More generally, a fairly large heuristic area: D.C. Programming  
exists  $\rightarrow$  we might revisit it, depending on importance.

~~Here's another example~~ For several examples of use of  
CCP in Stats & ML.

(a) "The Concave-Convex Procedure" — Yuille & Rangarajan

(b)

Note: This method is often a useful baseline to beat  
when dealing with nonconvex problems!

This repeated model-bound minimization is called  
Concave-Concave programming!

(9)

→ ~~CHAD~~ (one-member of a class of methods called Majorize-Minimize, etc.)  
(We'll see some more of this later in class) → But really handy  
tool to have in repertoire.

# A few words about gradient-descent methods:

Say  $\min f(x)$  s.t.  $x \in \mathbb{R}^n$   
•  $x$  s.t.  $\nabla f(x) \neq 0$ ;  
• Look at  $x_2 = x - \alpha \nabla f(x)$

# Methods without explicit derivatives

$$\min_{x \in \mathbb{R}^n} f(x) = f(x_1, \dots, x_n)$$

~~Very~~ at each iteration optimize over just 1 coordinate  
(or max if possible):

Called coordinate descent or Block-CG

# If  $f(x) = f(u, v)$ ,  $u \in \mathbb{R}^n$ ,  $v \in \mathbb{R}^m$

$$\text{then } u_{k+1} \leftarrow \arg\min_u f(u, v_k)$$

$$v_{k+1} \leftarrow \arg\min_v f(u_{k+1}, v)$$

Observe that  $\{f(u_k, v_k)\}$  is  $\downarrow$ .

# The two block case of alternating minimization is special →  
we will return to it later in the course.

5

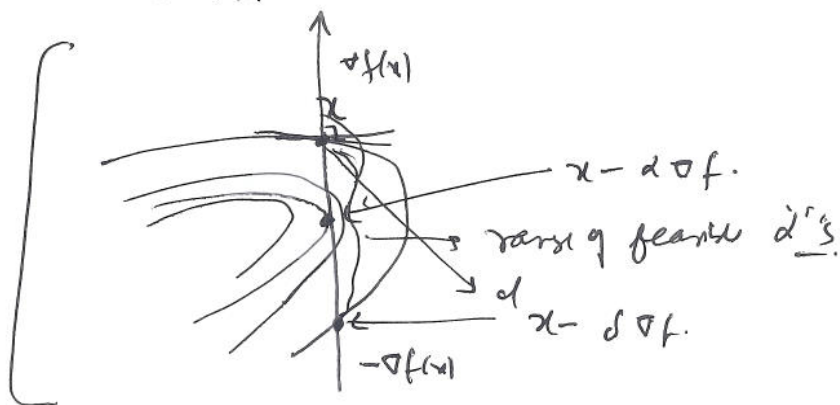
# # Back to <sup>descent</sup> gradient-methods:

for small enough  $\alpha > 0$   
we have descent

$$x \leftarrow x - \alpha \nabla f$$

more generally,  $x \leftarrow x + \alpha d$   $\leftarrow$  descent direction

This picture is  
important.



Code: while  $\neg$ done

$$x_{\text{new}} \leftarrow x_{\text{old}} + \alpha d_{\text{old}}$$

Update diff, done somehow.

end.

#  $x_{\text{new}}$  chosen to ensure  $\underline{f(x_{\text{new}}) < f(x_{\text{old}})}$  (Unless already stationary)

#  $d_{\text{old}}$ : descent direct.  $\underline{\langle \nabla f(x_{\text{old}}), d_{\text{old}} \rangle < 0}$

# Read "Numerical Optimization" by Nocedal & Wright to learn much more about  
different strategies for selecting  $\alpha$  and  $d$

# How about breaking the rules a little bit?

# Do NOT insist on  $f(x_{\text{new}}) < f(x_{\text{old}})$

# allows more flex in choosing step sizes.

# But why? maybe empirical speedups!

Observe:  $m(x, \alpha) = f(x) + \langle \nabla f(x), \alpha \rangle + \frac{\alpha^2}{2} \|\nabla^2 f(x)\|$  was ① model

How about the model.

$$m(x, \alpha) = f(x) + \langle \nabla f(x), \alpha \rangle + \frac{1}{2} \|\alpha\|^2$$

But now optimizing over  $\alpha$  ~~is not~~

6

- $x_{k+1} = x_k - \alpha S_k \nabla f(x_k)$  : Newton, Quasi-Newton
- ~~What if~~ ~~the~~ ~~choice~~ ~~is~~ : How to select  $S_k$  efficiently... (LBRs, etc.)
- Let us instead choose Set  $\alpha = 1$ ;  $S_k = \beta I$ ;

Find a good  $\beta$ !

$$S_k \approx [\nabla^2 f(x_k)]^{-1};$$

$$H \cdot \Delta x = \Delta g$$

$$H^{-1} \Delta g = \Delta x$$

$$\min_{\beta} \|\beta \Delta g - \Delta x\|^2 : \quad \beta^2 \langle \Delta g, \Delta g \rangle - 2\beta \langle \Delta g, \Delta x \rangle$$

$$\Rightarrow \beta = \frac{\langle \Delta g, \Delta x \rangle}{\|\Delta g\|^2} \quad !$$

alternatively, we can search for

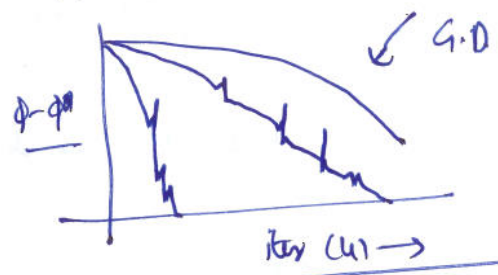
$$\|\beta \Delta x - \Delta g\|^2 \Rightarrow \beta = \frac{\langle \Delta x, \Delta g \rangle}{\|\Delta x\|^2} \Rightarrow \beta^{-1} = \frac{\|\Delta x\|^2}{\langle \Delta x, \Delta g \rangle} \quad !$$

step is  $\frac{1}{\beta}$

→ These steps are known as "Barzilai-Borwein" step sizes!

→ Empirically remarkable! "Barzilai & Borwein": Two-point step size gradient methods. IMA J. Num. Anal.; 1988.

Lead to picture such as :



# Let's ~~start~~ ~~now~~ begin now a a more careful study of theoretically evaluating performance of methods. But always, for a real problem, first compare on a computer to get a better idea! There is just a guideline.

7

$$O\left(\frac{1}{\epsilon}\right) \cdot \frac{1}{2^k} \approx 0 \quad \Rightarrow \quad \log k = -2^k \log 2 \quad (\log k = 2^k \log 2)$$

$$2^k \approx \frac{\log k}{\log 2}$$

To judge methods / their performance, we'll ~~use~~ could use:

- analytical complexity:
- Arithmetic complexity:

What does it mean to solve a problem:

Cur: obtain a point  $\bar{x}$  st  $f(\bar{x}) \leq f^* + \epsilon$ .

When substituting 'depends on  $\log\left(\frac{1}{\epsilon}\right)$  or  $\log\left(\log\left(\frac{1}{\epsilon}\right)\right)$ ,  $\log\left(\frac{1}{\epsilon}\right)$ '

Alternatively:  $\|\nabla f(\bar{x})\| \leq \epsilon$ . maybe meaningful.

or even  $\|x_{k+1} - (x_k - \alpha_k \nabla f(x_k))\|$  and so on.

## # Model of Computation:

# Types of oracles and examples of these oracles.

Deterministic [ Zeroth \_\_\_\_\_  
First \_\_\_\_\_  
Second \_\_\_\_\_

Stochastic [

More generally:

oracles with error.

that does not satisfy  $\mathbb{E}[q] = 0$ .

(Biased oracles)

# # Class of functions:

8

$C_L^{k,p}(\mathcal{X})$ :  $k$  times cont. diff func.

where  $p$ th deriv is Lipschitz cont with const  $L$  i.e.

$$\|f^{(p)}(x) - f^{(p)}(y)\| \leq L\|x - y\|, \quad \forall x, y \in \mathcal{X}.$$

Useful Lemma (HW):  $f \in C_L^{1,1}(\mathcal{X})$  iff  $\|f'(x)\| \leq L \quad \forall x \in \mathcal{X}$ .

Usage: What is l.c. of  $f(x) = x^T A x$ ? ;  $f(x) = \log(1+e^x)$ .

Notes:  $-\log x \notin C_L^{1,1}(\mathbb{R}_+)$  for any  $L$  |  $C_L^{1,1} \equiv C_L^{1,1}$   
 $\max(0, 1-x) \in C_1^{0,1}(\mathbb{R})$ .

Let us finally analyze

# Theorem: Let  $f \in C_{\text{cvx}}, C^1(\mathbb{R}^n)$ ; Let  $\alpha_k = 1/L$ .

Let  $x_{k+1} = x_k - \alpha_k \nabla f(x_k)$ ; then

$$f(x_k) - f(x^*) \leq \frac{2L \|x_0 - x^*\|^2}{k+1} = O(1/k)$$

i.e.  $O(1/k)$  iter to get  $\epsilon'$  acc.