

ML & Optimization Lecture 12

- Newton's Method
- Stochastic Gradient Descent

Newton's Method

Uses knowledge of the curvature to choose the update direction

$$f(x) \approx f(x_t) + \underbrace{\langle \nabla f(x_t), x - x_t \rangle + \frac{1}{2} (x - x_t)^T \nabla^2 f(x_t) (x - x_t)}_{\text{Taylor expansion}}$$

$$x_{t+1} = \underset{x}{\operatorname{argmin}}$$

By Fermat's condition we have x_{t+1} satisfies

$$0 = \nabla f(x_t) + \nabla^2 f(x_t) (x_{t+1} - x_t)$$

$$\Rightarrow x_{t+1} = x_t - [\nabla^2 f(x_t)]^{-1} \nabla f(x_t)$$

"Unguarded Newton's method"

We generally take our search direction

$$g_t = -[\nabla^2 f(x_t)]^{-1} \nabla f(x_t)$$

and judiciously choose our step-size α_t , then take

$$x_{t+1} = x_t + \alpha_t g_t$$

e.g. backtracking line search

This is called the guarded Newton's method

There are two phases of this alg:

- "damped phase": suboptimality goes down by a constant factor at each iteration: (here $\alpha_t < 1$)

$$f(x_t) - f(x_*) \leq \epsilon \quad \text{after } t = O(\log(\frac{1}{\epsilon}))$$

- "purely Newton phase": ($\alpha_t = 1$, BTLIS does nothing) at each iteration, the number of digits of accuracy doubles!

$$f(x_t) - f(x_*) \leq \epsilon \quad \text{after } t = O(\log \log(\frac{1}{\epsilon}))$$

Ex Consider an ill-conditioned linear system

Solving $x_* = \operatorname{argmin}_x \frac{1}{2} \|Ax - b\|_2^2$

$$f(x) = \frac{1}{2} \|Ax - b\|_2^2 \Rightarrow \nabla^2 f(x) = A^T A$$

and

$$\lambda_{\min}(A^T A) \cdot I \preceq A^T A \preceq \lambda_{\max}(A^T A) \cdot I$$

so

$$\kappa_f = \frac{\lambda_{\max}(A^T A)}{\lambda_{\min}(A^T A)}$$

Assume $\kappa_f = e^C$ where C is large.

Opt 1

Use GD with step size $\frac{1}{\beta} = \frac{1}{\lambda_{\max}(A^T A)} \ll 1$

$$f(x_T) - f(x_*) \leq \epsilon \quad \text{when} \quad T = O\left(K \log\left(\frac{1}{\epsilon}\right)\right) \gg 1$$

Opt 2

Use Newton's method with unit step size

$$x_1 = x_0 - (A^T A)^{-1} A^T (Ax_0 - b)$$

$$= (A^T A)^{-1} A^T b$$

$$= x_{\text{OLS}}$$

$$= \operatorname{argmin}_x \frac{1}{2} \|Ax - b\|_2^2$$

One iteration of
Newton's method
gives the exact
solution

Takeaway: we converge faster if
we account for curvature

Drawbacks to Newton's Method

- need to form $\nabla^2 f(x_t) \in \mathbb{R}^{d \times d}$
and solve a linear system

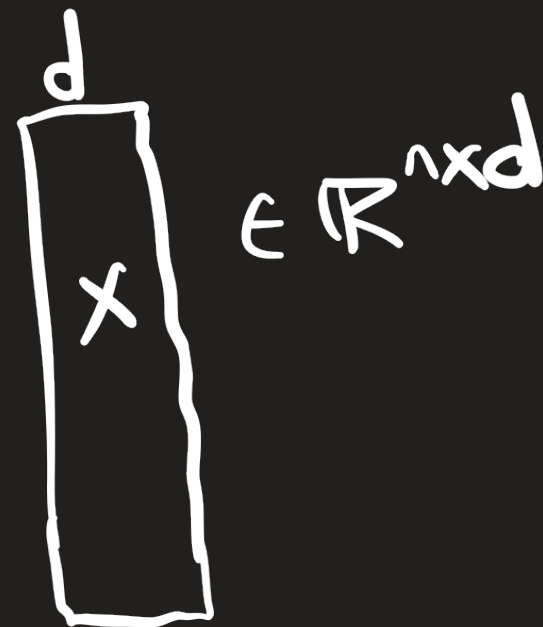
In general, this costs $O(d^2)$ memory and
 $O(d^3)$ operations

so when d is large, this is not feasible

GD can be expensive

Ex GD for solving a linear system

$$f(\omega) = \frac{1}{2} \|X\omega - y\|_2^2 = \sum_{i=1}^n \frac{1}{2} (x_i^T \omega - y_i)^2$$



$$\begin{aligned}\nabla f(\omega) &= X^T (X\omega - y) \\ &= X^T X \omega - X^T y\end{aligned}$$

Two approaches

1) can store $X^T X \in \mathbb{R}^{d \times d}$
then cost of computing
 $\nabla f(\omega)$ is $O(d^2)$

2) can't store $X^T X$
cost of computing $\nabla f(\omega)$
is $O(nd)$

Usually $n \gg d$, so both costs are $O(d^2)$ so if $d \gg 1$
then GD is expensive

Stochastic Gradient Descent

We want an alg cheaper per iteration than GD, preferably $O(d)$ per iteration

Idea: notice that $\nabla f(\omega_t)$ is an approximate search direction

$$f(\omega) \approx f(\omega_t) + \langle \nabla f(\omega_t), \omega - \omega_t \rangle + \frac{1}{2\alpha_t} \|\omega - \omega_t\|_2^2$$

where

$$f(\omega) = \frac{1}{n} \sum_{i=1}^n \ell(\omega^T x_i, y_i)$$

$$\Rightarrow \nabla f(\omega) = \frac{1}{n} \sum_{i=1}^n x_i \ell'(\omega^T x_i, y_i)$$

$\nabla f(\omega)$ is already a sample average that approximates a population average.

Treat $\nabla f(\omega)$ itself as a population average, and approximate it with a sample average

- select a random subset $I \subset [n] = \{1, \dots, n\}$ of size k and form an approximation

$$g = \frac{1}{k} \sum_{i \in I} x_i \ell'(\omega^T x_i, y_i) \quad \text{--- stochastic gradient estimate}$$

if we choose I appropriately

$$\mathbb{E} g = \nabla f(\omega)$$

The set I is called a minibatch

Stochastic Gradient Descent

$$x_* = \operatorname{argmin}_x f(x)$$

Choose a batch size k

for $t=0, \dots, T-1$

sample $I \subseteq [n]$ of size k

$$g_t \leftarrow \frac{1}{k} \sum_{i \in I} \nabla \ell_i(x_t)$$

$$\ell_i(w) := \ell(w^T x_i, y_i)$$

$$x_{t+1} = x_t - \alpha_t g_t$$

return

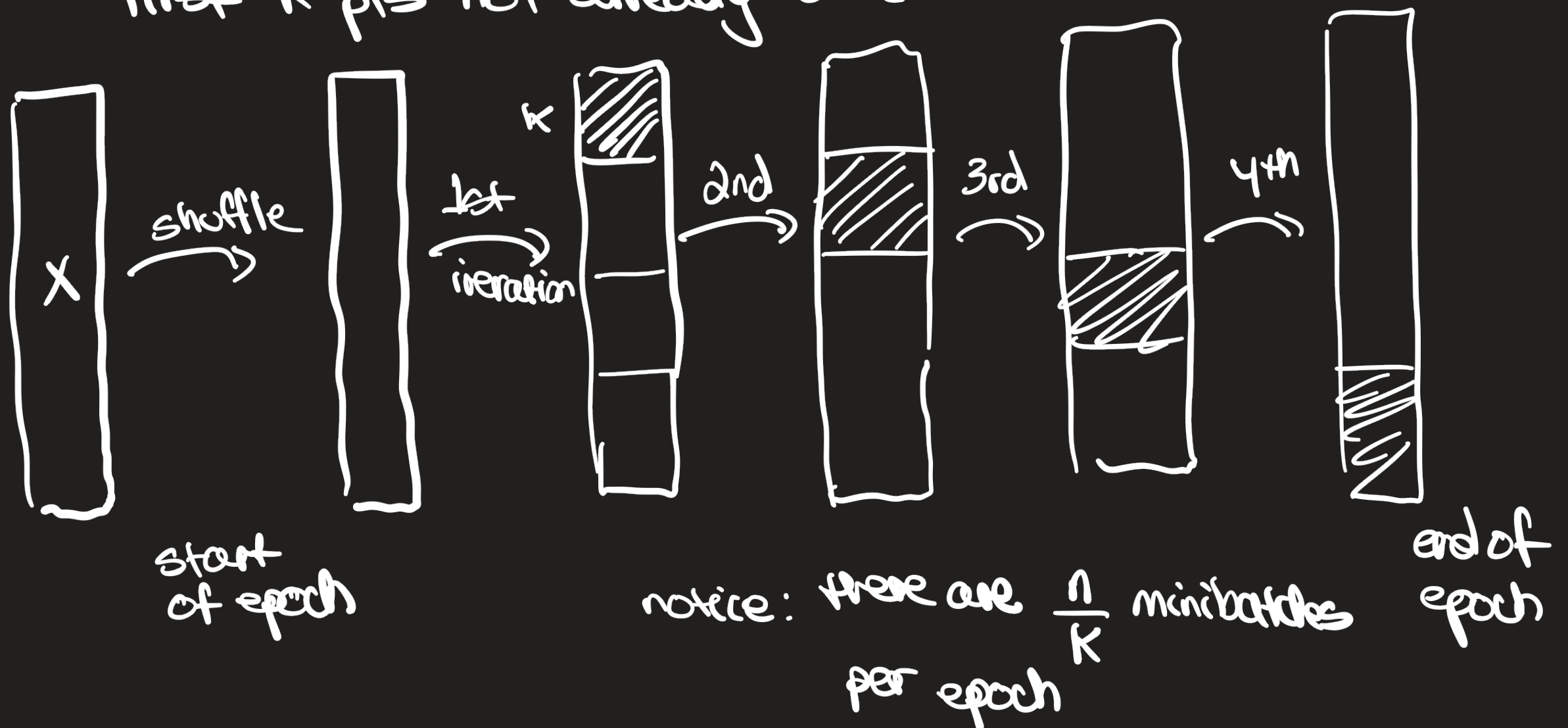
$$\hat{x} = \frac{1}{T} \sum_{i=1}^T x_i$$

Polyak average

In practice, to efficiently use all our data, we form our minibatches in the following way:

In each epoch:

- randomize the order of the training data
- everytime we sample a minibatch, use the first k pts not already used



Minibatch SGD

Inputs: ∇f oracle, α_t step sizes, k minibatch size, E epochs

Alg

$$z_0 = x_0$$

for $e=1, \dots, E$

$(X, y) \leftarrow$ random shuffling of itself

$$x_0 = z_{e-1}$$

for $t=0, \dots, \frac{n}{k}-1$

$I \leftarrow$ first k unused samples in this epoch

$$g_t \leftarrow \frac{1}{k} \sum_{i \in I} \nabla \ell_i(x_t)$$

$$x_{t+1} \leftarrow x_t - \alpha g_t$$

end

$$z_e = x_{\frac{n}{k}}$$

end

return z_E

Convergence rate of SGD (for nonconvex opt prob)

Due to nonconvexity and fact SGD is not a descent method, we measure convergence with

$$\min_{t=0, \dots, T-1} \mathbb{E} \|\nabla f(x_t)\|_2^2 \leq \epsilon$$

Assumptions

- f bounded below
- f is β -smooth

- assume $\mathbb{E} g_t = \nabla f(x_t)$

- $\mathbb{E} \|g_t\|_2^2 \leq \sigma^2$ (could relax to $\mathbb{E} \|\nabla f(x_t) - g_t\|_2^2 \leq \sigma^2$)

Idea: bound

$$\sum_{t=0}^{T-1} \alpha_t \mathbb{E} \|\nabla f(x_t)\|_2^2 \leq B$$

which implies

$$\min_{t=0, \dots, T-1} \mathbb{E} \|\nabla f(x_t)\|_2^2 \leq \frac{B}{\sum_{t=0}^{T-1} \alpha_t}$$

Prf

Recall that b/c f is β -smooth,

$$f(y) \leq f(x) + \langle \nabla f(x), y-x \rangle + \frac{\beta}{2} \|x-y\|_2^2$$

$$\begin{aligned} \Rightarrow f(x_{t+1}) &\leq f(x_t) + \langle \nabla f(x_t), -\alpha_t g_t \rangle + \frac{\beta}{2} \|-\alpha_t g_t\|_2^2 \\ &= f(x_t) - \alpha_t \langle \nabla f(x_t), g_t \rangle + \frac{\alpha_t^2 \beta}{2} \|g_t\|_2^2 \end{aligned}$$

Take expectations (b/c x_t, x_{t+1}, g_t are random vectors)

$$\mathbb{E} f(x_{t+1}) \leq \mathbb{E} f(x_t) - \alpha_t \mathbb{E} \|\nabla f(x_t)\|_2^2 + \frac{\alpha_t^2 \beta \sigma^2}{2}$$

this implies

$$\alpha_t \mathbb{E} \|\nabla f(x_t)\|_2^2 \leq \mathbb{E} [f(x_t) - f(x_{t+1})] + \frac{\alpha_t^2 \beta \sigma^2}{2}$$

so

$$\sum_{t=0}^{T-1} \alpha_t \mathbb{E} \|\nabla f(x_t)\|_2^2 \leq \sum_{t=0}^{T-1} \mathbb{E} [f(x_t) - f(x_{t+1})]$$

telescoping sum $\nearrow$ $+ \frac{\beta \sigma^2}{2} \sum_{t=0}^{T-1} \alpha_t^2$

$$= f(x_0) - \mathbb{E} f(x_T) + \frac{\beta \sigma^2}{2} \sum_{t=0}^{T-1} \alpha_t^2$$

$$\leq f(x_0) + \frac{\beta \sigma^2}{2} \sum_{t=0}^{T-1} \alpha_t^2$$

So finally

$$\underbrace{\min_{t=0, \dots, T-1} \mathbb{E} \|\nabla f(x_t)\|_2^2}_{=: \text{gap}(T)} \leq \frac{f(x_0)}{\sum_{t=0}^{T-1} \alpha_t} + \frac{\beta \sigma^2}{2} \frac{\sum_{t=0}^{T-1} \alpha_t^2}{\sum_{t=0}^{T-1} \alpha_t}$$

Choices of step-sizes:

$$\begin{aligned} - \alpha_t = \alpha, \text{ then } \text{gap}(T) &= O\left(\frac{f(x_0)}{T\alpha} + \frac{\beta \sigma^2}{2} \left(\frac{T\alpha^2}{T\alpha}\right)\right) \\ &= O\left(\frac{f(x_0)}{T} + \sigma^2 \alpha\right) \end{aligned}$$

$$- \alpha_t = \frac{\alpha}{t+1}, \text{ then } \sum_{t=0}^{T-1} \alpha_t = \sum_{t=0}^{T-1} \frac{\alpha}{t+1} = O(\ln T)$$

$$\text{and } \sum_{t=0}^{T-1} \alpha_t^2 = \sum_{t=0}^{T-1} \frac{\alpha^2}{(t+1)^2} = O(1) \Rightarrow \text{gap}(T) = O\left(\frac{1}{\ln T}\right)$$

$$- \alpha_t = \frac{\alpha}{\sqrt{t+1}} \quad \text{then} \quad \sum_{t=0}^{T-1} \alpha_t = \sum_{t=0}^{T-1} \frac{\alpha}{\sqrt{t+1}} = O(\sqrt{T})$$

$$\text{and} \quad \sum_{t=0}^{T-1} \alpha_t^2 = O(\ln T), \quad \text{so}$$

$$\text{gap}(T) = O\left(\frac{\ln T}{\sqrt{T}}\right)$$

If f is μ -strongly convex with const step size α

$$\mathbb{E} \|x_T - x_*\|_2^2 \leq (1 - 2\alpha\mu)^T \|x_0 - x_*\|_2^2 + \frac{\alpha\sigma^2}{2\mu}$$



after we stall in making progress on our training set, half the step size will allow us to get closer