

ML & Opt Lecture 10

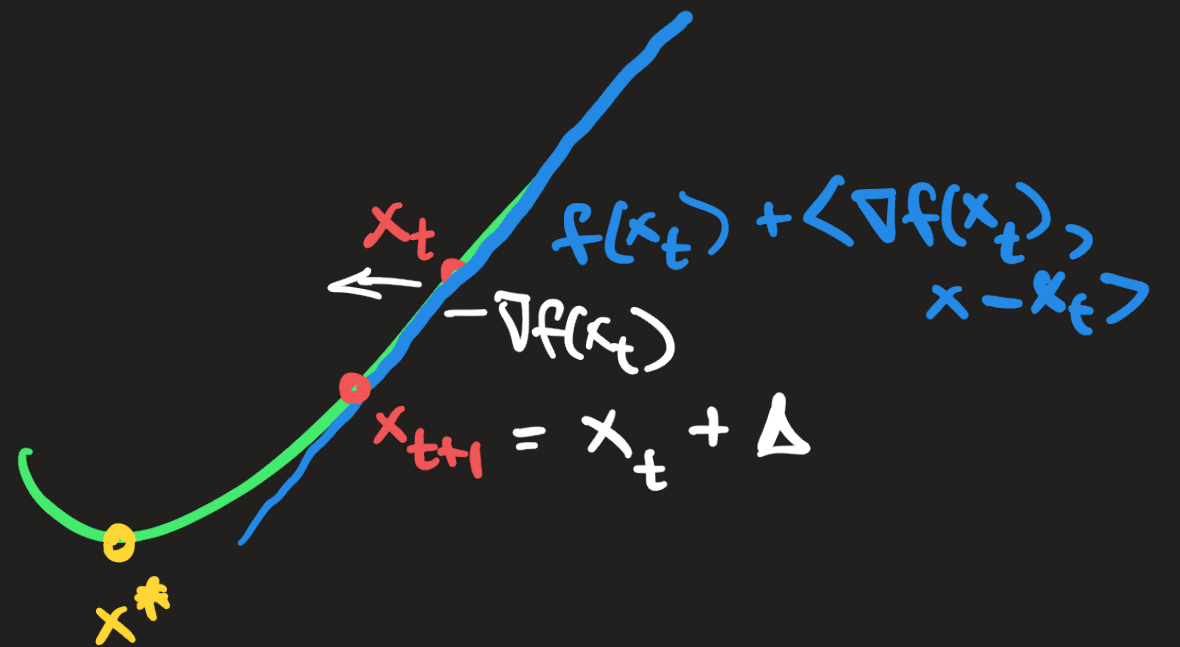
- Gradient Descent & Projected Gradient Descent
- Subgradient Descent & " Subgradient Descent
- Example

Gradient Descent

f - convex objective (smooth)

want solve

$$x_* \in \operatorname{argmin}_{x \in \mathbb{R}^d} f(x)$$



The basic idea is to note that by T.S. expansion argument:

$$\begin{aligned} f(x_{t+1}) &\approx f(x_t) + \langle \nabla f(x_t), x_{t+1} - x_t \rangle \\ &= f(x_t) + \langle \nabla f(x_t), \Delta \rangle \end{aligned}$$

since $\Delta = -\alpha \nabla f(x_t)$ pts in the neg grad dir,

$$\begin{aligned} f(x_{t+1}) &\approx f(x_t) - \alpha \langle \nabla f(x_t), \nabla f(x_t) \rangle \\ &= f(x_t) - \alpha \|\nabla f(x_t)\|_2^2 \end{aligned}$$

$< f(x_t)$, when α is small

Gradient Descent Alg

For solving

$$(U) \min_{x \in \mathbb{R}^d} f(x) \quad \text{where } f \text{ is smooth \& convex}$$

Input

∇f — gradient oracle

T — iterations

x_0 — starting guess

$\alpha_1, \dots, \alpha_T$ — step sizes for each iteration

Output

$x_1, \dots, x_T \in \mathbb{R}^d$ estimates of x^*

Alg

For $t = 0$ to $T-1$

$$x_{t+1} = x_t - \alpha_{t+1} \nabla f(x_t)$$

For now we want to know GD converges under some reasonable assumptions of f , that tell us:

- how to choose our stepsizes
- how many steps T to take to reach **approximate** convergence

ϵ -suboptimality

→ our focus for convex optim

Three notions of approximate convergence

1) $f(x_t) - f(x_*) < \epsilon$ approx conv in function value

2) $\|x_t - x_*\|_2^2 < \epsilon$ " " " parameter

3) $\|\nabla f(x_t)\|_2 < \epsilon$ approx conv of gradient to zero



↑
 ϵ -stationarity

usual convergence condition for nonconv optim

A reasonable assumption for GD to work is that f is β -smooth:

$$\|\nabla f(x) - \nabla f(y)\|_2 \leq \beta \|x - y\|_2,$$

i.e. ∇f is Lipschitz-continuous

Ex of a function that is β -smooth

$$f(x) = \|Ax - b\|_2^2 \quad \nabla f(x) = A^T(Ax - b)$$

$$\|\nabla f(x) - \nabla f(y)\|_2 = \|A^T A(x - y)\|_2 \leq \|A^T A\|_2 \|x - y\|_2$$

so f is β -smooth where $\beta = \|A^T A\|_2$

Ex of a function that is not β -smooth for any β

$$f(x) = e^x \quad \nabla f(x) = e^x$$

is there a β s.t. $|e^x - e^y| \leq \beta |x - y|$?

No, b/c there exists $z \in [x, y]$ s.t. $|e^x - e^y| = |e^z(x - y)| \geq e^x |x - y|$
so for any β I can find x & $y = x + 1$ so that $\beta < e^x$, so $|e^x - e^y| > \beta |x - y|$

Consequences of β -smoothness

$$1) \quad \|\nabla f(x)\|_2 \leq \beta \|x - x^*\|_2 \quad \left(\begin{array}{l} \text{take } y = x^* \text{ in the def} \\ \text{of } \beta\text{-smoothness \& note} \\ \nabla f(x^*) = 0 \end{array} \right)$$

2) This implies a concrete bound on how good the approximation

$$f(y) \approx f(x) + \nabla f(x)^T (y - x) \text{ is.}$$

Precisely,

$$|f(y) - [f(x) + \nabla f(x)^T (y - x)]| \leq \frac{\beta}{2} \|x - y\|_2^2$$

Chap 3 of Bubeck's Convex Optimization

Now assume f is β -smooth. Take $\alpha_t = \frac{1}{\beta}$

Descent Lemma with this choice, GD ensures that

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

Prf

$$f(x_{t+1}) = f(x_t - \alpha_t \nabla f(x_t))$$

$$\begin{aligned} \Rightarrow f(x_{t+1}) - \left[f(x_t) + \left\langle \nabla f(x_t), -\frac{1}{\beta} \nabla f(x_t) \right\rangle \right] \\ = f(x_{t+1}) - \left[f(x_t) - \frac{1}{\beta} \|\nabla f(x_t)\|_2^2 \right] \end{aligned}$$

$$\leq \frac{\beta}{2} \|x_{t+1} - x_t\|_2^2 = \frac{\beta}{2} \left\| \frac{1}{\beta} \nabla f(x_t) \right\|_2^2$$

$$= \frac{1}{2\beta} \|\nabla f(x_t)\|_2^2$$

Thus,

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

$\Rightarrow$

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

as claimed.

Fact If f is convex and β -smooth, then gradient descent with $\alpha_t = \frac{1}{\beta}$ guarantees that

$$f(x_T) - f(x_*) \leq \frac{2\beta \|x_0 - x_*\|_2^2}{T}$$

We say GrD takes $T = \mathcal{O}\left(\frac{\beta}{\epsilon}\right)$ iterations to get to an ϵ -suboptimal minimizer (iteration complexity)

Alternatively, GrD has convergence rate $\mathcal{O}\left(\frac{\beta}{T}\right)$

what about constrained optimization?

Projected Gradient Descent

$$(C) \min_C f(x)$$

Inputs

∇f - gradient oracle

T - iterations

x_0 - initial pt

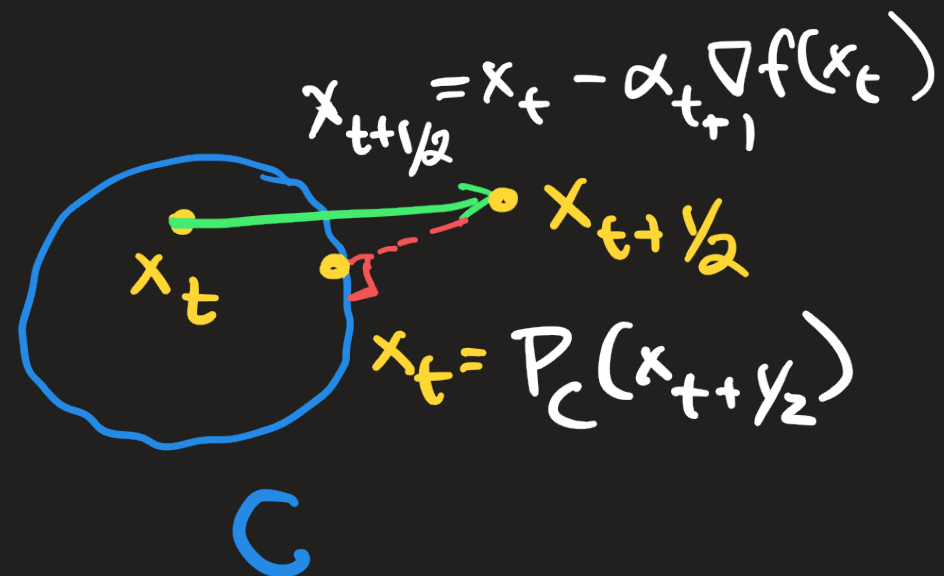
$\alpha_1, \dots, \alpha_T$ - step sizes

P_C - projection operator onto C

Alg

for t in 0 to $T-1$:

$$x_{t+1} = P_C(x_t - \alpha_{t+1} \nabla f(x_t))$$



What about nonsmooth, constrained convex optimization?

Subgradient Descent Alg

f - convex function

$$x^* = \operatorname{argmin}_{x \in C} f(x)$$

Inputs

∂f - oracle for a subgradient

T - # of iterations

x_0 - starting pt

$\alpha_1, \dots, \alpha_T$ - stepsizes

P_C - projector onto C

Alg

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

$g_t \leftarrow$ a vector in $\partial f(x_t)$

$$x_{t+1} \leftarrow P_C(x_t - \alpha_t g_t)$$

Return

$$\hat{x} = \frac{1}{T} \sum_{t=1}^T x_t \quad (\text{Polyak average})$$

Output

$$\hat{x} \approx x^*$$

Guarantees

$$f(\hat{x}) - f(x^*) \leq O\left(\frac{\log T}{T}\right)$$

