

ML and Optimization Lecture 13

- Drawbacks of GD

- stepsize choice : assumption f is Lipschitz, and we know its Lipschitz constant

- one soln (GD) : backtracking line search

- the curvature of our functions : suboptimal estimates from using Lipschitzianity of gradients

- use Hessian information : Newton's method

- the scalability of GD : generically, if have n datapoints in $\mathbb{R}^d$ each cost of GD is $O(nd)$!

- use estimates of the gradient : SGD $\frac{O(d)}{\text{cost per iteration}}$

Stepsize choice via Line search

Recall for GD that $-\nabla f(x)$ points towards the minimizer

$$f(x^*) \geq f(x) + \langle \nabla f(x), x^* - x \rangle \quad -\nabla f(x)$$

$$\Rightarrow \langle -\nabla f(x), x^* - x \rangle \geq 0$$



So far we've been choosing to move to

$$\hat{x} = x - \alpha \nabla f(x)$$

where stepsize α is chosen based on a priori knowledge of f being β -smooth ($\|\nabla f(x) - \nabla f(y)\|_2 \leq \beta \|x - y\|_2$ for all x, y): $\alpha = \frac{1}{\beta}$

Line search chooses step-sizes in an adaptive (but expensive) way to decrease the total cost of optimization.

Idea: minimize along the ray $x - \alpha \nabla f(x)$ the function value

$$\alpha_t = \operatorname{argmin}_{\alpha \geq 0} f(x_t - \alpha \nabla f(x_t))$$

This called exact line search and α_t is called the Curry step length. Not typically used, but can be useful for specific problems.

More generally we use backtracking line search.

Idea of backtracking line search:

T.S. expansion w/
remainder

$$f(x) = f(x_t) + \langle \nabla f(x_t), x - x_t \rangle + r(x - x_t), \text{ where } r(\Delta) = o(\|\Delta\|) \text{ as } \Delta \rightarrow 0$$

so take $x = x_t + \alpha d_t$ where d_t is a descent direction:

$$f(x_t + \alpha d_t) = f(x_t) + \alpha \underbrace{\langle \nabla f(x_t), d_t \rangle}_{\text{negative}} + r(\alpha d_t)$$

By the definition of $o(\|\Delta\|)$, for any $\delta > 0$, there is an $\varepsilon > 0$ such that

$$|r(\alpha d_t)| < \delta \|\alpha d_t\|_2 = \alpha \delta \|d_t\|_2. \text{ when } \alpha < \varepsilon$$

Take $\delta = (c-1) \frac{\langle \nabla f(x_t), d_t \rangle}{\|d_t\|_2}$, where $c \in (0, 1)$, so $\delta > 0$, then

$$\begin{aligned} \alpha \langle \nabla f(x_t), d_t \rangle + r(\alpha d_t) &< \alpha \langle \nabla f(x_t), d_t \rangle + \\ &\quad (c-1) \alpha \langle \nabla f(x_t), d_t \rangle \\ &= c \alpha \langle \nabla f(x_t), d_t \rangle \end{aligned}$$

so we can guarantee that if α is small enough, we get decrease:

$$f(x_t + \alpha d_t) < f(x_t) + c \alpha \langle \nabla f(x_t), d_t \rangle$$

- Backtracking line search "backtracks" to find a small enough α to ensure decrease.

Backtracking Line Search

Given: x_t , d_t - descent direction (i.e. a vector s.t. $\langle d_t, -\nabla f(x_t) \rangle \geq 0$)

$c \in (0, 1)$ - amount of requested decrease

$\beta \in (0, 1)$ - backtracking parameter

Return: α s.t.

$$f(x_t + \alpha d_t) < f(x_t) + c\alpha \langle \nabla f(x_t), d_t \rangle$$

Alg

$\alpha \leftarrow 1$

while $f(x_t + \alpha d_t) \geq f(x_t) + c\alpha \langle \nabla f(x_t), d_t \rangle$

$\alpha \leftarrow \beta \alpha$

return α

The cost of BTLS depends on:

- 1) The cost of evaluating f
- 2) How much decrease you ask for (how large c is)
- 3) The aggressiveness of your backtracking (how small β is)

Idea: balance aggressiveness & decrease asked for to make maximal progress per call to BTLS

(Boyd & Vandenberghe. Convex Optimization)

General prescriptions:

$$c \in [0.01, 0.3]$$

$$\beta \in [0.5, 0.8]$$

Gradient Descent w/ Exact Line Search

(on μ -strongly convex & β -smooth convex functions)

$$f(x_T) - f(x_*) \leq \left(1 - \frac{1}{K}\right)^{T-1} (f(x_1) - f(x_*))$$

where $K = \frac{\beta}{\mu}$ is the condition number of the convex function f

β = an upper bound on the eigenvalues of $\nabla^2 f$
on the entire domain

μ = a lower bound on the eigenvalues of $\nabla^2 f$
on the entire domain

$$K \leq \frac{\beta}{\mu} \iff \mu I \preceq \nabla^2 f(x) \preceq \beta I$$

Strong convexity If

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

then we say f is μ -strongly convex

Clear that $\mu \leq \beta$, so $K = \frac{\beta}{\mu} \geq 1$

The convergence bound says

$$f(x_T) - f(x_*) \leq c^{T-1} (f(x_1) - f(x_*))$$

where $c = 1 - \frac{1}{K}$, so $K \geq 1$ implies

steepest descent (GD + line search) converges fast.

Newton's Method

We saw that steepest descent works well when $\kappa = \frac{L}{\mu}$ is small, which means f looks ^{locally} like a quadratic w/ the same amount of curvature in every direction.

What if we don't have this: e.g. consider

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

where $\nabla^2 f(x) = A^T A$
is ill-conditioned

Then steepest descent has poor convergence. Notice we solve this in one step with

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

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

$$= x_0 - \nabla^2 f(x_0)^{-1} \nabla f(x_0)$$

Note that $-\nabla^2 f(x_t)^{-1} \nabla f(x_t) = d_t$ is a descent direction.

Why?

We know $\nabla^2 f(x_t) \succeq 0$ so its eigenvalues are ≥ 0
this means $\nabla^2 f(x_t)^{-1} \succeq 0$ as well, so

$$\langle d_t, -\nabla f(x_t) \rangle = \langle \nabla^2 f(x_t)^{-1} \nabla f(x_t), \nabla f(x_t) \rangle \geq 0$$



Another view of Newton's method: we explicitly model f 's curvature

$$\begin{aligned} f(x) &\approx f(x_t) + \langle \nabla f(x_t), x - x_t \rangle \\ &\quad + \frac{1}{2} (x - x_t)^T \nabla^2 f(x_t) (x - x_t) \\ &= m(x) \end{aligned}$$

Since the approximation $m(x)$ is convex we can approximately minimize f by setting

$$\begin{aligned} \nabla m(x)_{t+1} = 0 &\iff \nabla f(x_t) + \nabla^2 f(x_t)(x_{t+1} - x_t) = 0 \\ &\iff x_{t+1} = x_t - \nabla^2 f(x_t)^{-1} \nabla f(x_t) \end{aligned}$$

This suggests we can use $\nabla^2 f(x_t)^{-1} \nabla f(x_t)$ as our search direction.

As before we can line search to choose how far to move in that direction

Damped/guarded Newton's method

Input: x_t , oracles for f , ∇f , $\nabla^2 f$

Output: x_T

Alg

For $t = 1, \dots, T-1$

$$d_t = -\nabla^2 f(x_t)^{-1} \nabla f(x_t)$$

$$\alpha_t \leftarrow \text{BTLS}$$

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

What is the benefit of using curvature info?

There are two phases of the algorithm:

— damped phase: suboptimality goes down by
($\alpha_t < 1$) a constant factor at every iteration

in this phase: $f(x_t) - f(x_*) \leq \varepsilon$ $O(\log(\frac{1}{\varepsilon}))$

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

in this

phase: $f(x_t) - f(x_*) \leq \varepsilon$ $O(\log \log(\frac{1}{\varepsilon}))$

see Boyd & Vandenberghe for details and proof

Drawbacks:

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

In general this costs $\mathcal{O}(d^2)$ memory and
 $\mathcal{O}(d^3)$ operations

so when d is large (say 100k) this is not
feasible.