

This is a draft.

1 Review of Last Lecture

1.1 Steepest descent method

The (unnormalized) gradient of Steepest descent method is defined as

$$g_{sd} = -\nabla f(x)$$

where the normalized gradient

$$g_{nsd} = \arg \min_v \|\nabla f(x) - v\|_g$$

and the $\|\cdot\|_g$ denote the dual norm.

For a positive matrix A , we define matrix norm w.r.t. A as

$$\|x\|_A^2 = x^T A x$$

which can be used in the dual norm.

2 Newton's Method

2.1 definition

The gradient of Newton's Method is defined as

$$x_{nt} = -(\nabla^2 f(x))^{-1} \nabla f(x)$$

which actually is a special case of steepest descent if we use matrix norm w.r.t. $\nabla^2 f(x)$, i.e., the Hessian matrix of f at this point.

Notice that Newton's method sometimes may not converge¹.

The intuition of Newton's Method is to approximate the function f by a quadratic function $\hat{f}$, then compute the minimize of $\hat{f}$ to get the new point.

Example 1 (Application on Quadratic Function) Suppose f is a quadratic function ,

$$\hat{f}(x+v) = f(x) + \nabla f(x)^T v + \frac{1}{2} v^T \nabla^2 f(x) v$$

Then

$$\nabla \hat{f}(x+v) = 0 \Rightarrow \nabla f(x) + \nabla^2 f(x) v = 0; v = -(\nabla^2 f(x))^{-1} \nabla f(x)$$

Which means that we only need 1 iteration to get optimal.

Remark 2 Newton's method is a descent method, since the inner product $\langle x_{nt}, \nabla f(x) \rangle = -\nabla f(x)^T \nabla^2 f(x)^{-1} \nabla f(x)$ is guaranteed to be negative.

¹See http://en.wikipedia.org/wiki/Newton's_method#Failure_analysis.

2.2 Affine Invariance

Consider $f: \mathbb{R}^n \rightarrow \mathbb{R}$, and an invertible linear transformation $T \in \mathbb{R}^{n \times n}$.

$$\bar{f}(y) = f(Ty)$$

Where T is a $n \times n$ invertible matrix.

It is easy to verify the following

$$\nabla \bar{f}(y) = T^T \nabla f(x); \quad \nabla^2 \bar{f}(x) = T^T \nabla^2 f(x) T;$$

and the step size of $\bar{f}$ at point x is

$$y_{nt} = (T^T \nabla^2 f(x) T)^{-1} (T^T \nabla f(x)) = T^{-1} x_{nt}$$

where $x = Ty$.

Remark 3 Since computing the term $(\nabla^2 f(x))^{-1}$ is usually a difficult job, an improvement of Newton's Method called Quasi-Newton is proposed.

3 Convergence Analysis

3.1 Convergence of Newton's method

Based on the following assumptions,

$$mI \preceq \nabla^2 f \preceq MI; \quad \|\nabla f(x) - \nabla f(y)\| \leq L\|x - y\|$$

the implement of Newton's method consists of the following two phases.

Phase 1. $\|\nabla f\| > \epsilon$, then $f(x^{k+1}) - f(x^k) \geq \epsilon$, where $\epsilon = 2m^2 = L^2$ and $0 < \epsilon < m^2 = L$.

Phase 2. Otherwise,

$$\frac{L}{2m^2} \|\nabla f(x_{k+1})\| \leq \left(\frac{L}{2m^2} \|\nabla f(x_k)\| \right)^2$$

Recall that in last class we mentioned that for strongly convex function f , $f(x) - p \leq \frac{1}{2m} \|\nabla f(x)\|^2$, which means that $f(x)$ is close to optimal when its gradient is small.

Then the number of iterations is

$$\frac{f(x_0) - p}{\epsilon} + \log \log \left(\frac{0}{\epsilon} \right)$$

Self-concordance is

$$\frac{f(x_0) - p}{\square} + \log \log \left(\frac{1}{\square} \right)$$

Where $\square$ is something irrelevant with

Note that the convergency result was for Newton's direction combined with backtrack line search (See notes in last class), not for the plain Newton step written in the board of this class.

Remark 4 In real application, since phase 2 performs large better than phase 1, one can first use other gradient descent method, and then use Newton's method when $\|\nabla f\|$ is sufficiently small.

3.2 Convergence of Gradient Descent (Constant Step Size)

Assumption 5 (Lipschitz gradient)

$$\| \nabla f(x) - \nabla f(y) \| \leq L \|x - y\|$$

Suppose that $t_k = (2 - \epsilon)/L$, where L is the coefficient in the Lipschitz assumption, then we have

Lemma 6 (Descent Lemma)

$$f(x_{k+1}) - f(x_k) \leq \nabla f(x_k)^T (x_{k+1} - x_k) + \frac{L}{2} \|x_{k+1} - x_k\|^2$$

Proposition 7 The following four conditions are equivalent.²

1. Lipschitz gradient.
2. Descent Lemma
3. Coercivity.

$$(\nabla f(x) - \nabla f(y))^T (x - y) \geq M \|x - y\|^2$$

4.

$$\| \nabla^2 f \| \leq L$$

Then we prove the convergence of f_{x_k} , with the number of iterations in $O(\cdot)$.

Proof: Using the lemma above,

$$\begin{aligned} f(x_{k+1}) - f(x_k) &\leq \nabla f(x_k)^T (x_{k+1} - x_k) + \frac{L}{2} \|x_{k+1} - x_k\|^2 \\ &= \|\nabla f(x_k)\|^2 \left(t_k + \frac{L}{2} t_k^2 \right) \\ &= \|\nabla f(x_k)\|^2 \left(t_k + t_k \frac{2}{2} \right) \\ &= \frac{2}{2} t_k \|\nabla f(x_k)\|^2 \\ &= \frac{2}{2} \|\nabla f(x_k)\|^2 \end{aligned}$$

By convexity of f ,

$$\begin{aligned} f(x_k) - f(x_0) &\leq \langle \nabla f(x_k), x_k - x_0 \rangle \\ &\leq \|\nabla f(x_k)\| \|x_k - x_0\| \\ &\leq R \|x_k - x_0\| \end{aligned}$$

where R is the radius defined as $R = \max_{f(x) = f(x_0)} \|x - x_0\|$.

²The lemma is recommended to prove in the order 1 $\Rightarrow$ 2 $\Rightarrow$ 3 $\Rightarrow$ 4 $\Rightarrow$ 1, left as homework.

Let $\epsilon_k = f(x_k) - f^*$, then

$$\epsilon_k - \epsilon_{k+1} = \frac{1}{2} \| \nabla f(x_k) \|^2 = \frac{1}{2} \frac{\epsilon_k^2}{R^2}$$

Therefore

$$\frac{1}{\epsilon_{k+1}} - \frac{1}{\epsilon_k} = \frac{k}{\epsilon_k \epsilon_{k+1}} = \frac{k}{\frac{1}{2} \frac{\epsilon_k^2}{R^2}} = \frac{2}{\epsilon_k} \Rightarrow \frac{1}{\epsilon_k} = \frac{1}{\epsilon_0} + \frac{k}{2R^2}$$

where $\epsilon_0 = f(x_0) - f^* \leq LR$.

So, we get that k is the order in $O(\frac{1}{\epsilon})$ to satisfy $\epsilon_k \leq \epsilon$. □

3.3 Lower Bounds for First Order Methods

Assume that

f is strongly convex with parameter L , i.e. $\nabla^2 f \succeq LI$.

The initial point $x_0 = 0$.

(First order)

$$x_k \in \text{Span}\{x_1, \dots, x_{k-1}; \nabla f(x_1), \dots, \nabla f(x_{k-1})\}$$

Then we show that the convergence rate of $f(x) = \frac{1}{2} x^T A x - b^T x$ is in $(\frac{L}{1} \log \frac{1}{\epsilon})$, where

$$A = A_0 + LI = \frac{L}{4} I + \begin{bmatrix} 2 & 1 & 0 & 0 & 0 \\ 1 & 2 & 1 & 0 & 0 \\ 0 & 1 & 2 & 1 & 0 \\ 0 & 0 & 1 & 2 & 0 \\ \vdots & \vdots & \vdots & \vdots & \ddots \\ 0 & 0 & 0 & 0 & 2 \end{bmatrix} + LI; \quad b = e_1$$

and dimension $d \geq 1$.

Claim 8 Suppose $Ax = b$ and $x = (u_1, \dots, u_k, \dots, u_d)$, then

$$u_i = \left(\frac{1}{1 + \frac{L}{4}} \right)^i$$

Key observations,

$x_0 = 0$, and $\nabla f(x_0) = b$.

$x_1 \in \text{Span}\{b\}$.

$x_2 \in \text{Span}\{Ab, b\}$.

$x_3 \in \text{Span}\{A^2b, Ab, b\}$.

⋮

$x_k \in \text{Span}\{A^{k-1}b; \dots; Ab; b\}$. (Krylov subspace)

Also notice that x_k must be in the form that only the first k terms are non-zero.

Therefore,

$$\|Ax - x\|^2 = \sum_{j=k+1}^d u_j^2 \|x\|^2 = u_{k+1}^2 \|x\|^2$$

where $u = (p - 1)/(p + 1)$. Hence

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

which implies that $k = \lceil \log \frac{1}{\epsilon} \rceil$ to satisfy $\|f(x_k) - f(x)\| < \epsilon$.

4 Conjugate Gradient

Goal: minimize

$$x^T A x - b^T x$$

given that $A \succ 0$.

Suppose that $K = \text{Span}\{b; Ab; \dots; A^{k-1}b\}$, and $\{v_0; \dots; v_{k-1}\}$ is a basis of K .

Then the question is equivalent to

Finding the best vector x in K that minimizes

$$\|x - \sum_i \alpha_i v_i\|_A$$

Observation.

If v_i and v_j are A -orthogonal to each other, i.e., $v_i^T A v_j = 0$; $i \neq j$, then

$$\left\| x - \sum_i \alpha_i v_i \right\|_A^2 = \sum_i (\alpha_i^2 v_i^T A v_i - 2 \alpha_i b^T v_i) + \|x\|_A^2$$

(all the cross terms $v_i^T A v_j$ ($i \neq j$) disappear)

which implies that $\alpha_i = b^T v_i / v_i^T A v_i$.

KEY IDEA, we can construct the basis $\{v_0; \dots; v_{k-1}\}$ which are A -orthogonal to each other efficiently step by step, i.e.,

$$\text{Span}\{v_0; \dots; v_i\} = \text{Span}\{b; Ab; \dots; A^i b\} \stackrel{\text{denote}}{=} K_i$$

Claim 9 $A v_{i-1}$ is A -orthogonal to $v_0; \dots; v_{i-2}$.

One-line proof: for $j \in [i-2]$, $A v_j \in K_{j+1}$, which is A -orthogonal to v_{i-1} . Hence $v_{i-1}^T A A v_j = 0 = (A v_{i-1})^T A v_j$, which prove the claim. $\square$

Therefore in each iteration, we only need to A -orthogonalize $A v_{i-1}$, v_{i-1} and v_{i-2} to obtain v_i . Since by construction v_i can write as a linear combination of $A v_{i-1}; v_{i-1}; v_{i-2}$, and by Claim 9 and induction $A v_{i-1}; v_{i-1}; v_{i-2}$ are A -orthogonal to $v_0; \dots; v_{i-3}$, we have v_i is A -orthogonal to $v_0; \dots; v_{i-3}$, which proves the Key idea.

References

- [1] Cover T M. On determining the irrationality of the mean of a random variable. Ann. Statist, 1973, 1(5): 862-871.