

Lecture 5: CS395T Numerical Optimization for Graphics and AI — Line Search

Qixing Huang
The University of Texas at Austin
huangqx@cs.utexas.edu

1 Disclaimer

This note is adapted from Section 3 of

- *Numerical Optimization* by Jorge Nocedal and Stephen J. Wright. Springer series in operations research and financial engineering. Springer, New York, NY, 2. ed. edition, (2006)

2 Introduction

Each iteration of a line search method computes a search direction $\mathbf{p}_k$ and then decides how far to move along that direction. The iteration is given by

$$\mathbf{x}_{k+1} = \mathbf{x}_k + \alpha_k \mathbf{p}_k, \quad (1)$$

where the positive scalar α_k is called the step length. The success of a line search method depends on effective choices of both the direction $\mathbf{p}_k$ and the step length α_k . Most line search algorithms require $\mathbf{p}_k$ to be a descent direction one for which $\mathbf{p}_k^T \nabla f(\mathbf{x}_k) < 0$ —because this property guarantees that the function f can be reduced along this direction, as discussed in the previous chapter. Moreover, the search direction often has the form

$$\mathbf{p}_k = -B_k^{-1} \nabla f(\mathbf{x}_k), \quad (2)$$

where B_k is a symmetric and non-singular matrix. In the steepest descent method B_k is simply the identity matrix I , while in Newton's method B_k is the exact Hessian $\nabla^2 f(\mathbf{x}_k)$. In quasi-Newton methods, B_k is an approximation to the Hessian that is updated at every iteration by means of a low-rank formula. When $\mathbf{p}_k$ is defined by (2) and B_k is positive definite, we have

$$\mathbf{p}_k^T \nabla f(\mathbf{x}_k) = -\nabla f(\mathbf{x}_k)^T B_k^{-1} \nabla f(\mathbf{x}_k) < 0,$$

and therefore $\mathbf{p}_k$ is a descent direction. In the next few lectures we study how to choose the matrix B_k , or more generally, how to compute the search direction. We now give careful consideration to the choice of the step-length parameter α_k .

3 Step Length

In computing the step length α_k , we face a tradeoff. We would like to choose α_k to give a substantial reduction of f , but at the same time, we do not want to spend too much time making the choice. The ideal choice would be the global minimizer of

$$\phi(\alpha) = f(\mathbf{x}_k + \alpha \mathbf{p}_k), \alpha > 0, \quad (3)$$

but in general, it is too expensive to identify this value. To find even a local minimizer of ϕ to moderate precision generally requires too many evaluations of the objective function f and possibly the gradient ∇f . More practical strategies perform an inexact line search to identify a step length that achieves adequate reductions in f at minimal cost. Typical line search algorithms try out a sequence of candidate values for α , stopping to accept one of these values when certain conditions are satisfied. The line search is done in two stages: A bracketing phase finds an interval containing desirable step lengths, and a bisection or interpolation phase computes a good step length within this interval.

How to terminate a line-search is critical. We now discuss various termination conditions for the line search algorithm and show that effective step lengths need not lie near minimizers of the univariate function $\phi(\alpha)$ defined above. A simple condition we could impose on α_k is that it provide a reduction in f , i.e., $f(\mathbf{x}_k + \alpha_k \mathbf{p}_k) < f(\mathbf{x}_k)$. However, this may not be appropriate. For example, consider minimizing $f(x) = (x+1)^2$. We can let $x_k = \frac{5}{k}$. It is clear that the value of the objective function always goes down. However, it does not go to the minimizer. The difficulty is that we do not have sufficient reduction in f , a concept we discuss next.

3.1 The Wolfe Conditions

A popular inexact line search condition stipulates that α_k should first of all give sufficient decrease in the objective function f , as measured by the following inequality:

$$f(\mathbf{x}_k + \alpha \mathbf{p}_k) \leq f(\mathbf{x}_k) + c_1 \alpha \nabla f(\mathbf{x}_k)^T \mathbf{p}_k, \quad (4)$$

for some constant $c_1 \in (0, 1)$. In other words, if you move far, then you should have more reduction. This implicitly favors small step-sizes. In practice, c_1 is chosen to be small, i.e., $c_1 = 10^{-4}$.

The sufficient decrease condition is not enough by itself to ensure that the algorithm makes reasonable progress. To rule out unacceptably short steps we introduce a second requirement, called the curvature condition, which requires α_k to satisfy

$$\nabla f(\mathbf{x}_k + \alpha_k \mathbf{p}_k)^T \mathbf{p}_k \geq c_2 \nabla f(\mathbf{x}_k)^T \mathbf{p}_k, \quad (5)$$

for some constant $c_2 \in (c_1, 1)$, where c_1 is the constant from (4). Note that the left-hand-side is simply the derivative $\phi'(\alpha_k)$, so the curvature condition ensures that the slope of $\phi(\alpha_k)$ is greater than c_2 times the gradient $\phi'(0)$. This makes sense because if the slope $\phi(\alpha)$ is strongly negative, we have an indication that we can reduce f significantly by moving further along the chosen direction. On the other hand, if the slope is only slightly negative or even positive, it is a sign that we cannot expect much more decrease in f in this direction, so it might make sense to terminate the line search. Typical values of c_2 are 0.9 when the search direction $\mathbf{p}_k$ is chosen by a Newton or quasi-Newton method, and 0.1 when $\mathbf{p}_k$ is obtained first order methods.

The sufficient decrease and curvature conditions are known collectively as the *Wolfe conditions*. We restate them here for future reference:

$$f(\mathbf{x}_k + \alpha_k \mathbf{p}_k) \leq f(\mathbf{x}_k) + c_1 \alpha_k \nabla f(\mathbf{x}_k)^T \mathbf{p}_k, \quad (6)$$

$$\nabla f(\mathbf{x}_k + \alpha_k \mathbf{p}_k)^T \mathbf{p}_k \geq c_2 \nabla f(\mathbf{x}_k)^T \mathbf{p}_k, \quad (7)$$

with $0 < c_1 < c_2 < 1$.

Lemma 3.1. *Suppose that $f : \mathbb{R}^n \rightarrow \mathbb{R}$ is continuously differentiable. Let $\mathbf{p}_k$ be a decent direction at $\mathbf{x}_k$, and assume that f is bounded below along the ray $\{\mathbf{x}_k + \alpha \mathbf{p}_k | \alpha > 0\}$. Then if $0 < c_1 < c_2 < 1$, there exist intervals of step lengths satisfying the Wolfe conditions.*

The proof is straight-forward using mean-value theorem.

3.2 Sufficient Decrease and Backtracking

We have mentioned that the sufficient decrease condition alone is not sufficient to ensure that the algorithm makes reasonable progress along the given search direction. However, if the line search algorithm chooses its candidate step lengths appropriately, by using a so-called backtracking approach, we can dispense with the extra condition use just the sufficient decrease condition to terminate the line search procedure. In its most basic form, backtracking proceeds as follows:

- Choose $\bar{\alpha} > 0, \rho, c \in (0, 1)$; set $\alpha \leftarrow \bar{\alpha}$;
- **repeat** until $f(\mathbf{x}_k + \alpha \mathbf{p}_k) \leq f(\mathbf{x}_k) + c\alpha \nabla f(\mathbf{x}_k)^T \mathbf{p}_k$.
- $\alpha \leftarrow \rho \alpha$;
- **end**.

4 Convergence of Line Search Methods

To obtain global convergence, we must not only have well-chosen step lengths but also well-chosen search directions $\mathbf{p}_k$. We discuss requirements on the search direction in this section, focusing on one key property: the angle θ_k between $\mathbf{p}_k$ and the steepest descent direction $-\nabla f(\mathbf{x}_k)$, defined by

$$\cos(\theta_k) = -\frac{\nabla f(\mathbf{x}_k)^T \mathbf{p}_k}{\|\nabla f(\mathbf{x}_k)\| \|\mathbf{p}_k\|}. \quad (8)$$

The following theorem, due to Zoutendijk, has far-reaching consequences. It shows, for example, that the steepest descent method is globally convergent. For other algorithms it describes how far $\mathbf{p}_k$ can deviate from the steepest descent direction and still give rise to a globally convergent iteration. Various line search termination conditions can be used to establish this result, but for concreteness we will consider only the Wolfe conditions.

Theorem 4.1. *Consider any iteration of the form (1), where $\mathbf{p}_k$ is a descent direction and α_k satisfies the Wolfe conditions. Suppose that f is bounded below in $\mathbb{R}^n$ and that f is continuously differentiable in an open set $\mathcal{N}$ containing the level set $\mathcal{L} := \{\mathbf{x} : f(\mathbf{x}) \leq f(\mathbf{x}_0)\}$, where $\mathbf{x}_0$ is the starting point of the iteration. Assume also that the gradient ∇f is Lipschitz continuous on $\mathcal{N}$, that is, there exists a constant $L > 0$ such that*

$$\|\nabla f(\mathbf{x}) - \nabla f(\bar{\mathbf{x}})\| \leq L \|\mathbf{x} - \bar{\mathbf{x}}\|, \quad \forall \mathbf{x}, \bar{\mathbf{x}} \in \mathcal{N}. \quad (9)$$

Then

$$\sum_{k \geq 0} \cos^2(\theta_k) \|\nabla f(\mathbf{x}_k)\|^2 < \infty.$$

Proof: From the line search conditions, we have that

$$(\nabla f(\mathbf{x}_{k+1}) - \nabla f(\mathbf{x}_k))^T \mathbf{p}_k \geq (c_2 - 1) \nabla f(\mathbf{x}_k)^T \mathbf{p}_k,$$

while the Lipschitz condition implies that

$$(\nabla f_{k+1} - \nabla f_k)^T \mathbf{p}_k \leq \alpha_k L \|\mathbf{p}_k\|^2.$$

By combining these two relations, we obtain

$$\alpha_k \geq \frac{c_2 - 1}{L} \frac{\nabla f(\mathbf{x}_k)^T \mathbf{p}_k}{\|\mathbf{p}_k\|^2}.$$

By substituting this inequality into the first Wolfe condition, we obtain

$$f(\mathbf{x}_{k+1}) \leq f(\mathbf{x}_k) - c_1 \frac{1 - c_2}{L} \frac{(\nabla f(\mathbf{x}_k)^T \mathbf{p}_k)^2}{\|\mathbf{p}_k\|^2},$$

or in other words,

$$f(\mathbf{x}_{k+1}) \leq f(\mathbf{x}_k) - c \cos^2(\theta_k) \|\nabla f(\mathbf{x}_k)\|^2,$$

where $c = c_1(1 - c_2)/L$. We can conclude the proof by summing this expression over all indices less than or equal to k , we obtain

$$f_{k+1} \leq f_0 - c \sum_{j=0}^k \cos^2(\theta_j) \|\nabla f(\mathbf{x}_j)\|^2.$$

□

Steepest decent: $\lim_{k \rightarrow \infty} \|\nabla f(\mathbf{x}_k)\| = 0$.

Newton method:

$$\cos(\theta_k) \geq \frac{1}{M},$$

where $\|B_k\| \|B_k^{-1}\| \leq M$.

4.1 Convergence Rate

Convergence rate of steepest descent: Let us suppose that

$$f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T Q \mathbf{x} - \mathbf{b}^T \mathbf{x},$$

where Q is symmetric and positive definite. The gradient is given by $\nabla f(\mathbf{x}) = Q\mathbf{x} - \mathbf{b}$, and the minimizer $\mathbf{x}^*$ is the unique solution of the linear system $Q\mathbf{x} = \mathbf{b}$.

Let us compute the step length α_k that minimizes $f(\mathbf{x}_k - \alpha \nabla f(\mathbf{x}_k))$. By differentiating

$$f(\mathbf{x}_k - \alpha \mathbf{g}_k) = \frac{1}{2} (\mathbf{x}_k - \alpha \mathbf{g}_k)^T Q (\mathbf{x}_k - \alpha \mathbf{g}_k) - \mathbf{b}^T (\mathbf{x}_k - \alpha \mathbf{g}_k)$$

with respect to α , we obtain

$$\alpha = \frac{\nabla f(\mathbf{x}_k)^T \nabla f(\mathbf{x}_k)}{\nabla f(\mathbf{x}_k)^T Q \nabla f(\mathbf{x}_k)}.$$

If we use this exact minimizer α_k , the steepest descent iteration is given by

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \left(\frac{\nabla f(\mathbf{x}_k)^T \nabla f(\mathbf{x}_k)}{\nabla f(\mathbf{x}_k)^T Q \nabla f(\mathbf{x}_k)} \right) \nabla f(\mathbf{x}_k).$$

Theorem 4.2. *When the steepest descent method with exact line searches is applied, the error norm satisfies*

$$\|\mathbf{x}_{k+1} - \mathbf{x}^*\|_Q^2 \leq \left(\frac{\lambda_n - \lambda_1}{\lambda_n + \lambda_1} \right)^2 \|\mathbf{x}_k - \mathbf{x}^*\|_Q^2 \quad (10)$$

Note that for this function of interest,

$$\frac{1}{2} \|\mathbf{x} - \mathbf{x}^*\|_Q^2 = (\mathbf{x} - \mathbf{x}^*)^T Q (\mathbf{x} - \mathbf{x}^*) = f(\mathbf{x}) - f(\mathbf{x}^*).$$

The rate of convergence behavior of the steepest descent method is essentially the same on general non-linear objective functions. In the following result we assume that the step length is the global minimizer along the search direction.

Theorem 4.3. *Suppose that $f : \mathbb{R}^n \rightarrow \mathbb{R}$ is twice continuously differentiable, and that the iterates generated by the steepest descent method with exact line searches converge to a point $\mathbf{x}^*$ where the Hessian matrix $\nabla^2 f(\mathbf{x}^*)$ is positive definite. Then*

$$(f(\mathbf{x}_{k+1}) - f(\mathbf{x}^*)) \leq \left(\frac{\lambda_n - \lambda_1}{\lambda_n + \lambda_1} \right)^2 (f(\mathbf{x}_k) - f(\mathbf{x}^*)). \quad (11)$$

Newton method:

Theorem 4.4. *Suppose that f is twice differentiable and that the Hessian $\nabla^2 f(\mathbf{x})$ is Lipschitz continuous in a neighborhood of a solution $\mathbf{x}^*$ at which the sufficient conditions are satisfied. Consider the iteration $\mathbf{x}_{k+1} = \mathbf{x}_k + \mathbf{p}_k$, where $\mathbf{p}_k = -(\nabla^2 f(\mathbf{x}_k))^{-1} \nabla f(\mathbf{x}_k)$. Then*

- 1. if the starting point $\mathbf{x}_0$ is sufficiently close to $\mathbf{x}^*$, the sequence of iterates converge to $\mathbf{x}^*$;*
- 2. the rate of convergence of $\mathbf{x}_k$ is quadratic; and*
- 3. the sequence of gradient norms $\|\nabla f(\mathbf{x}_k)\|$ converges quadratically to zero.*