

Learning objectives.

1. Compute the *gradient* and *Hessian* of simple multivariate functions.
2. Apply *gradient descent*, *stochastic gradient descent*, and *SGD with momentum* from a formula sheet.
3. Know why *ADAM* is basically just a variant of GD.
4. Apply *Newton's method* by solving the Newton linear system.
5. Convert equality constraints into a quadratic-penalty objective and inequality constraints into a log-barrier objective.

The emphasis is deliberately plug-and-chug. At this stage, students should recognize general patterns and know how to execute these methods correctly. Convergence analysis comes later in the course.

1 Linear algebra review

1.1 Vectors

Vector $x \in \mathbb{R}^n$ lies in n -dimensional Euclidean space:

$$x = (x_1, \dots, x_n) := \begin{bmatrix} x_1 & x_2 & \cdots & x_n \end{bmatrix}^T = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix}.$$

Inner product

$$\begin{aligned} \langle a, x \rangle &:= a^T x = \begin{bmatrix} a_1 & a_2 & \cdots & a_n \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} \\ &= a_1 x_1 + a_2 x_2 + \cdots + a_n x_n \end{aligned}$$

Euclidean norm.

$$\|x\| = \sqrt{\langle x, x \rangle} = \sqrt{x^T x} = \left(\sum_{i=1}^n x_i^2 \right)^{1/2}.$$

Two important results for the Euclidean norm (not used today, but will be important for next lecture):

1. *Pythagorean theorem*:

$$x^T y = 0 \iff \|x + y\|^2 = \|x\|^2 + \|y\|^2.$$

2. *Cauchy-Schwarz inequality*:

$$|x^T y| \leq \|x\| \|y\|.$$

ℓ_p norms (for $p \geq 1$).

$$\|x\|_p = \left(\sum_{i=1}^n |x_i|^p \right)^{1/p}, \quad \|x\|_\infty = \max_{i \in \{1, \dots, n\}} |x_i|.$$

All of these norms satisfy the following properties:

1. $\|x\|_p \geq 0$,
2. $\|x\|_p = 0 \iff x = 0$,
3. $\|cx\|_p = |c| \|x\|_p$ for $c \in \mathbb{R}$,
4. $\|x + y\|_p \leq \|x\|_p + \|y\|_p$ (triangle inequality).

1.2 Matrices

Matrix $A \in \mathbb{R}^{m \times n}$ lies in an mn -dimensional Euclidean space

$$A_{m \times n} = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{bmatrix} = \begin{bmatrix} - & a_1^T & - \\ - & a_2^T & - \\ & \vdots & \\ - & a_m^T & - \end{bmatrix}.$$

Matrix-vector product

$$\begin{aligned}
 Ax = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} &= \begin{bmatrix} - & a_1^T & - \\ - & a_2^T & - \\ & \vdots & \\ - & a_m^T & - \end{bmatrix} \begin{bmatrix} | \\ x \\ | \end{bmatrix} \\
 &= \begin{bmatrix} a_1^T x \\ a_2^T x \\ \vdots \\ a_m^T x \end{bmatrix} = \begin{bmatrix} a_{11}x_1 + a_{12}x_2 + a_{13}x_3 \\ a_{21}x_1 + a_{22}x_2 + a_{23}x_3 \\ a_{31}x_1 + a_{32}x_2 + a_{33}x_3 \end{bmatrix}
 \end{aligned}$$

Matrix inversion. A square matrix A is invertible if there is a unique matrix A^{-1} satisfying $A^{-1}A = AA^{-1} = I$. In that case,

$$Ax = b \iff x = A^{-1}b.$$

Two-by-two matrix inversion (invertible iff $ad \neq bc$)

$$\begin{bmatrix} a & b \\ c & d \end{bmatrix}^{-1} = \frac{1}{ad - bc} \begin{bmatrix} d & -b \\ -c & a \end{bmatrix}.$$

Question 1. Why do they say to **never** form the inverse A^{-1} ?

Forward substitution (invertible iff A_{ii} invertible for all i)

$$\begin{bmatrix} A_{11} & & & \\ A_{21} & A_{22} & & \\ \vdots & \vdots & \ddots & \\ A_{n1} & A_{n2} & \cdots & A_{nn} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} = \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{bmatrix}$$

Special case: block diagonal (invertible iff A_{ii} invertible for all i)

$$\begin{bmatrix} A_{11} & & & \\ & A_{22} & & \\ & & \ddots & \\ & & & A_{nn} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} = \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{bmatrix}$$

1.3 Gradients and Hessians

Let $f(x_1, \dots, x_n)$ be real-valued. Then

$$\nabla f(x) = \begin{bmatrix} \partial f / \partial x_1 \\ \partial f / \partial x_2 \\ \vdots \\ \partial f / \partial x_n \end{bmatrix}.$$

Hessian is the Jacobian of the gradient (symmetric for the smooth functions considered here)

$$\nabla^2 f(x) =$$

2 Gradient descent

For $\min_x f(x)$, gradient descent with stepsize $\eta > 0$ is written

$$x_{k+1} = x_k - \eta \underbrace{\nabla f(x_k)}_{g_k}.$$

We commonly write g_k for the gradient evaluated at the k -th iteration.

Example 1. Gradient descent. Let

$$f(x) = x_1^4 + 2x_2^2 + \exp(-x_3) + \exp(2x_3).$$

1. Derive $\nabla f(x)$.
2. Perform one gradient descent iteration from $x_0 = (1, 2, 0)$ with stepsize $\eta = 1/2$.

3 Newton's method

For $\min_x f(x)$, Newton's method with stepsize $\eta > 0$ is written

$$x_{k+1} = x_k - \eta \underbrace{\nabla^2 f(x_k)^{-1}}_{B_k} \underbrace{\nabla f(x_k)}_{g_k}.$$

The matrix B_k is called the *preconditioner*. (Sometimes, this name also refers to B_k^{-1} .) It is intended to make GD converge faster.

The resulting method is a kind of “preconditioned GD.”

Example 2. Newton's method. Let

$$f(x) = x_1^4 + 2x_2^2 + \exp(-x_3) + \exp(2x_3).$$

1. Derive $\nabla^2 f(x)$.
2. Perform one **full** Newton iteration from $x_0 = (1, 2, 0)$.

4 Stochastic gradient descent

SGD is designed to optimize a finite-sum, expectation-like objective

$$f(x) = \frac{1}{m} \sum_{i=1}^m f_i(x) \approx \mathbb{E}[\mathbf{f}(x)].$$

Recall that GD computes the full-batch gradient

$$\nabla f(x_k) = \frac{1}{m} \sum_{i=1}^m \nabla f_i(x_k).$$

This is considered “slow.”

Question 2. Why?

SGD instead computes the *stochastic gradient*. In this lecture, we use random reshuffling: at the start of each epoch, shuffle the m components and iterate through them without replacement. With a single component,

$$x_{k+1} = x_k - \eta g_k \quad \text{where } g_k = \nabla f_i(x_k),$$

or with a *minibatch* B from the shuffled order,

$$g_k = \frac{1}{|B|} \sum_{i \in B} \nabla f_i(x_k), \quad |B| = N.$$

Exhausting the complete set of m samples this way is called an *epoch*.

Without qualifiers, we understand “SGD” to mean $N = 1$.

Example 3. Stochastic gradient descent. Let $f(x) = \frac{1}{2}(f_1(x) + f_2(x))$, where

$$f_1(x) = \frac{1}{2}(x - 1)^2, \quad f_2(x) = \frac{1}{2}(x - 3)^2.$$

Starting from $x_0 = 0$ with $\eta = 1/2$:

1. take one gradient-descent step;
2. take one epoch of SGD in the order 1, 2.

5 SGD with momentum

Let g_k denote the stochastic gradient used at iteration k . Momentum keeps a running average of past gradients, with parameter $\beta \in [0, 1)$

$$v_k = \beta v_{k-1} + (1 - \beta)g_k,$$

$$x_k = x_{k-1} - \eta v_k.$$

This is a classical momentum update. It is inspired by a closely-related scheme that was shown by Nesterov to achieve optimal convergence rate for smooth convex optimization.

Question 3. How to “turn off” momentum?

Example 4. Momentum. Let $f_1(x) = f_2(x) = \frac{1}{2}(x - 2)^2$. Starting from $x_0 = 0$ and $v_0 = 0$, perform one epoch of SGD with momentum using $\eta = 1/2$ and $\beta = 1/2$.

6 ADAM

Adam is the classic optimizer for neural-network training and was famously used to train GPT-3 and InstructGPT, the precursors to ChatGPT. Since 2024, matrix-aware optimizers such as Muon have become more common. Adam is effectively just a combination of the three ideas discussed so far: (a) preconditioning; (b) stochastic gradients; (c) momentum acceleration.

Let g_k be the stochastic gradient from before. Then ADAM is just

$$\begin{aligned}v_k &= \beta_1 v_{k-1} + (1 - \beta_1) g_k, \\x_k &= x_{k-1} - \eta_k B_k v_k,\end{aligned}$$

where the diagonal preconditioner B_k is constructed in a special way by accumulating the previous gradients

$$\begin{aligned}u_k &= \beta_2 u_{k-1} + (1 - \beta_2)(g_k \odot g_k), \\B_k &= \frac{1}{1 - \beta_1^k} \left[\text{diag} \left(\frac{u_k}{1 - \beta_2^k} \right)^{1/2} + \varepsilon I \right]^{-1}.\end{aligned}$$

Question 4. What is the difference between ADAM and...

- SGD+momentum?
- SGD?
- Classical gradient descent?

7 How to deal with constraints?

So far, all algorithms have been unconstrained.

What if we have constraints?

$$\min_{x_1, x_2} f(x_1, x_2) \text{ s.t. } x_1^2 + x_2^2 = 1.$$

Question 5. Why do we square the constraint?

Question 6. What does the parameter μ do?

Question 7. Why doesn't this “solve” constrained optimization?

8 What about inequality constraints?

What if we have *inequality* constraints?

$$\min_{x_1, x_2} f(x_1, x_2) \text{ s.t. } x_1^2 + x_2^2 \leq 1.$$

Question 8. What is the difference?

Question 9. What does the parameter μ do?

Question 10. Why doesn't this “solve” constrained optimization?

9 Convergence diagnostics

Common quantities to plot are

$$f(x_k), \quad f(x_k) - f^*, \quad \|\nabla f(x_k)\|, \quad \|x_k - x^*\|.$$

They answer different questions.

$f(x_k)$	Is the objective decreasing?
$f(x_k) - f^*$	How far is the objective value from optimal?
$\ \nabla f(x_k)\ $	Is the method approaching a stationary point?
$\ x_k - x^*\ $	Is the iterate approaching a known reference solution?

For the penalty and barrier formulations above, plot $\|\nabla F_\mu(x_k)\|$; the gradient of the original constrained objective need not vanish at a constrained optimum.

A flat plot of $f(x_k)$ alone is not evidence of failure. If f^* is large, $f(x_k)$ can look visually stagnant even while the objective gap and gradient norm decay rapidly.