

Machine Learning: Optimization 3

Hiroshi Shimodaira and Hao Tang

2025

Ver. 1.0

Topics

- Gradient descent
- Step size / learning rate
- Approximate solution
- Convergence rate
- Stochastic gradient descent (SGD)
- Mini-batch and epoch

Convexity on more points

If a function f is convex,

$$f(\alpha_1 \mathbf{x}_1 + \alpha_2 \mathbf{x}_2 + \alpha_3 \mathbf{x}_3) \leq \alpha_1 f(\mathbf{x}_1) + \alpha_2 f(\mathbf{x}_2) + \alpha_3 f(\mathbf{x}_3) \quad (1)$$

for $\alpha_1, \alpha_2, \alpha_3 \geq 0$ and $\alpha_1 + \alpha_2 + \alpha_3 = 1$.

Convexity on more points (*cont.*)

If a function f is convex,

$$f\left(\sum_{i=1}^N \alpha_i \mathbf{x}_i\right) \leq \sum_{i=1}^N \alpha_i f(\mathbf{x}_i) \quad (2)$$

for $\alpha_i \geq 0$ and $\sum_{i=1}^N \alpha_i = 1$.

Jensen's inequality

If a function f is convex,

$$f(\mathbb{E}[\mathbf{x}]) \leq \mathbb{E}[f(\mathbf{x})]. \quad (3)$$

$$f(\mathbb{E}_{\mathbf{x} \sim p(\mathbf{x})}[\mathbf{x}]) \leq \mathbb{E}_{\mathbf{x} \sim p(\mathbf{x})}[f(\mathbf{x})].$$

- For log loss

$$\text{NLL} = \sum_{i=1}^N \log \left(1 + \exp(-y_i \mathbf{w}^\top \phi(\mathbf{x}_i)) \right) \quad (4)$$

we cannot even solve $\nabla_{\mathbf{w}} \text{NLL} = \mathbf{0}$.
(i.e. no closed-form solution)

- How do we find the optimal solution?

Gradient descent

- Gradient descent is an iterative algorithm, consisting of the steps

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t \nabla \text{NLL}(\mathbf{w}_t). \quad (5)$$

- The variable $\eta_t > 0$ is called the **step size** (or learning rate), and can depend on t .

Gradient descent on log loss

- The log loss in the binary case

$$\text{NLL} = \sum_{i=1}^N \log \left(1 + \exp(-y_i \mathbf{w}^\top \mathbf{x}_i) \right). \quad (6)$$

- We have shown that NLL is convex in $\mathbf{w}$.

Gradient descent on log loss (*cont.*)

$$\frac{\partial \text{NLL}}{\partial \mathbf{w}} = \sum_{i=1}^N \frac{\exp(-y_i \mathbf{w}^\top \mathbf{x}_i)}{1 + \exp(-y_i \mathbf{w}^\top \mathbf{x}_i)} (-y_i \mathbf{x}_i) \quad (7)$$

$$= \sum_{i=1}^N \left(1 - \frac{1}{1 + \exp(-y_i \mathbf{w}^\top \mathbf{x}_i)} \right) (-y_i \mathbf{x}_i) \quad (8)$$

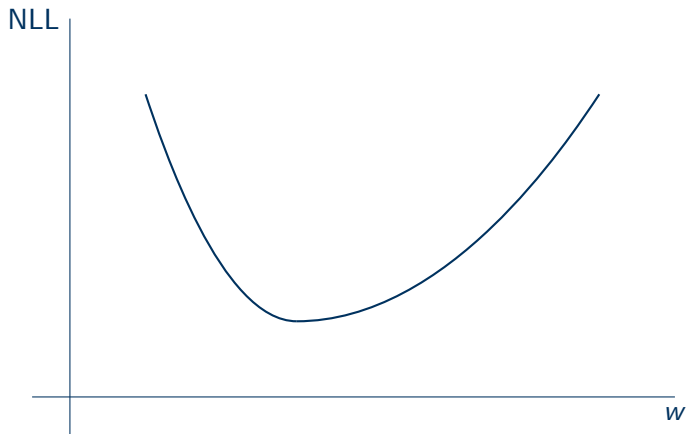
$$= \sum_{i=1}^N (1 - p(y_i | \mathbf{x}_i)) (-y_i \mathbf{x}_i) \quad (9)$$

Gradient descent on log loss (*cont.*)

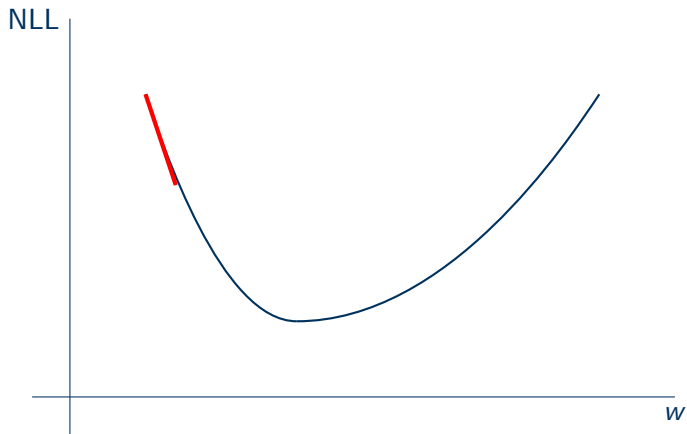
$$\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t \nabla \text{NLL}(\mathbf{w}_t) \quad (10)$$

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t \sum_{i=1}^N \left(1 - \frac{1}{1 + \exp(-y_i \mathbf{w}_t^\top \mathbf{x}_i)} \right) (-y_i \mathbf{x}_i) \quad (11)$$

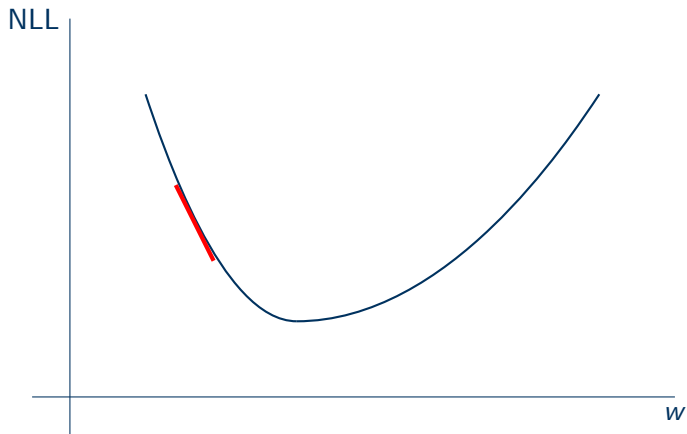
Gradient descent



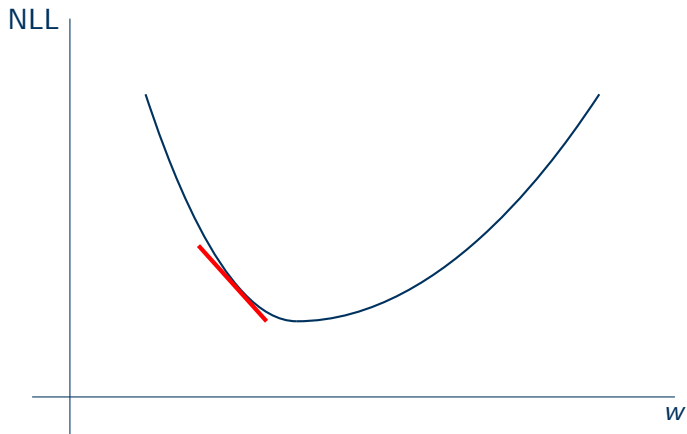
Gradient descent



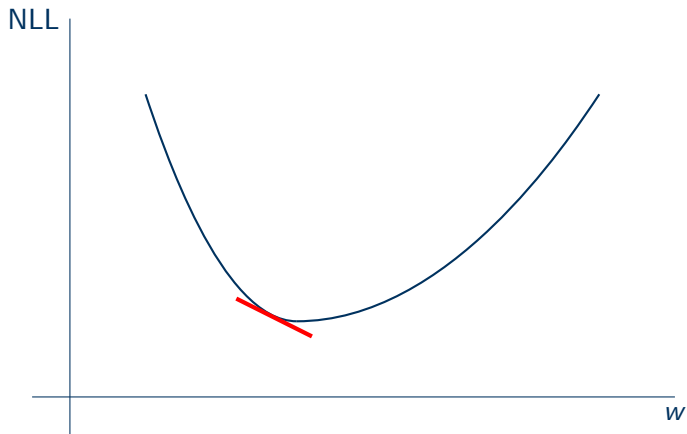
Gradient descent



Gradient descent



Gradient descent



Approximate solutions in optimisation

- We say that $\hat{\mathbf{x}}$ is an approximate solution of the minimiser $\mathbf{x}^*$ if, for a given $\epsilon > 0$,

$$f(\hat{\mathbf{x}}) - f(\mathbf{x}^*) < \epsilon. \quad (12)$$

- Note that it is close in function value, not close in the input.

Approximate solutions for iterative algorithms

- An iterative algorithm creates a sequence $\mathbf{x}_1, \dots, \mathbf{x}_t$.
- How many updates do we need to achieve an approximate solution?
- Given $\epsilon > 0$, how large does t needs to be to achieve

$$f(\mathbf{x}_t) - f(\mathbf{x}^*) < \epsilon? \tag{13}$$

- We want to express ϵ as a function of t .

Rate of convergence

$$\lim_{t \rightarrow \infty} \frac{|f(\mathbf{x}_{t+1}) - f(\mathbf{x}^*)|}{|f(\mathbf{x}_t) - f(\mathbf{x}^*)|^q} = \mu \quad (14)$$

The sequence $\{f(\mathbf{x}_t)\}$ is said to *converge* with order q to $f(\mathbf{x}^*)$ and with a *rate of convergence* μ

- For $q = 1$
 - If $\mu = 1$, *sublinear rate*
 - If $0 < \mu < 1$, *linear rate*
 - if $\mu = 0$, *superlinear rate*

Potential results

- Sublinear

- $f(\mathbf{x}_t) - f(\mathbf{x}^*) \leq \frac{c}{t^2}$

- Linear

- $f(\mathbf{x}_t) - f(\mathbf{x}^*) \leq cr^t$ for $0 < r < 1$

- Quadratic

- $f(\mathbf{x}_t) - f(\mathbf{x}^*) \leq cr^{2^t}$ for $0 < r < 1$

Potential results

- Sublinear

- $f(\mathbf{x}_t) - f(\mathbf{x}^*) \leq \frac{c}{t^2}$
- $\epsilon = O\left(\frac{1}{t^2}\right)$ or $t = O\left(\frac{1}{\sqrt{\epsilon}}\right)$

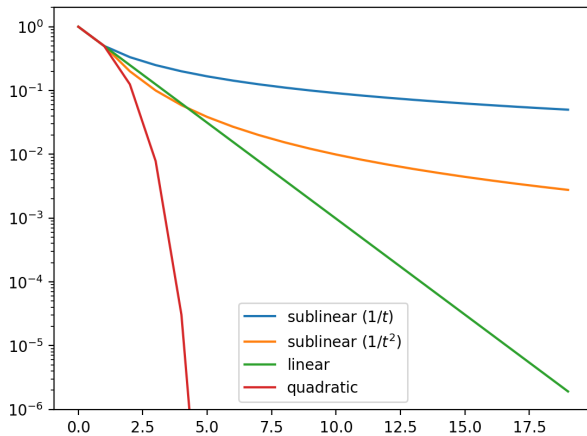
- Linear

- $f(\mathbf{x}_t) - f(\mathbf{x}^*) \leq cr^t$ for $0 < r < 1$
- $\epsilon = O(2^{-t})$ or $t = O(\log \frac{1}{\epsilon})$

- Quadratic

- $f(\mathbf{x}_t) - f(\mathbf{x}^*) \leq cr^{2^t}$ for $0 < r < 1$
- $\epsilon = O\left(2^{-2^t}\right)$ or $t = O(\log \log \frac{1}{\epsilon})$

Convergence rates



Gradient descent on mean-squared error

- The mean-squared error is

$$L = \|\mathbf{X}\mathbf{w} - \mathbf{y}\|_2^2. \quad (15)$$

- We have shown that L is convex in $\mathbf{w}$.
- We have shown that the optimal solution is $(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$.

Gradient descent on mean-squared error (*cont.*)

$$\nabla L = 2(\mathbf{X}^\top \mathbf{X} \mathbf{w} - \mathbf{X}^\top \mathbf{y}) \quad (16)$$

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t \nabla L(\mathbf{w}_t) = \mathbf{w}_t - \eta_t 2(\mathbf{X}^\top \mathbf{X} \mathbf{w}_t - \mathbf{X}^\top \mathbf{y}) \quad (17)$$

Runtime comparison for solving mean-squared error

- The runtime of $(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$ is $O(nd^2)$ or $O(d^3)$ (whichever dominates).

Runtime comparison for solving mean-squared error

- The runtime of $(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$ is $O(nd^2)$ or $O(d^3)$ (whichever dominates).

Runtime comparison for solving mean-squared error

- The runtime of $(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$ is $O(nd^2)$ or $O(d^3)$ (whichever dominates).
- The runtime of a single gradient step $\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t 2(\mathbf{X}^\top \mathbf{X} \mathbf{w}_t - \mathbf{X}^\top \mathbf{y})$ is $O(nd)$.

Convergence rates of gradient descent on mean-squared error

- If we run gradient descent on mean-squared error, we have

$$L(\mathbf{w}_t) - L(\mathbf{w}^*) = \frac{1}{2}(\mathbf{w}_0 - \mathbf{w}^*)^\top (\mathbf{I} - \eta \mathbf{H})^{2t} \mathbf{H}(\mathbf{w}_0 - \mathbf{w}^*) \quad (18)$$

where $\mathbf{H} = \mathbf{X}^\top \mathbf{X}$ is the Hessian of L .

- If we choose $\eta = \frac{1}{\lambda_{\max}}$, where $\lambda_{\max}$ is the largest eigenvalue of $\mathbf{H}$, the convergence is linear.

Stochastic gradient descent (SGD)

1. Sample $\mathbf{x}_t, \mathbf{y}_t$ from a data set S .
2. $\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t \nabla \ell(\mathbf{w}_t; \mathbf{x}_t, \mathbf{y}_t)$
 - Per sample L_2 loss $\ell(\mathbf{w}; \mathbf{x}, \mathbf{y}) = (\mathbf{w}^\top \mathbf{x}_t - y_t)^2$
 - Per sample log loss $\ell(\mathbf{w}; \mathbf{x}, \mathbf{y}) = \log(1 + \exp(-y_t \mathbf{w}^\top \mathbf{x}_t))$
3. Go to 1 until the solution is satisfactory.

Stochastic gradient descent (SGD) (*cont.*)

- $\nabla \ell(\mathbf{w}_t; \mathbf{x}_t, y_t)$ is now random, because $\mathbf{x}_t$ and y_t is random.
- The expectation

$$\mathbb{E}_{\mathbf{x}, y \sim U(S)}[\nabla \ell(\mathbf{w}; \mathbf{x}, y)] = \nabla L(\mathbf{w}) \quad (19)$$

where $U(S)$ is the uniform distribution over the samples in S .

Guarantee for stochastic gradient descent

- If we do SGD on an convex function,

$$\mathbb{E}_{\mathbf{x}, \mathbf{y} \sim U(S)}[L(\bar{\mathbf{w}}_t)] - L(\mathbf{w}^*) \leq \frac{\|\mathbf{w}_0 - \mathbf{w}^*\|_2^2}{2\eta t} + \frac{\eta B^2}{2}. \quad (20)$$

- $\|\nabla \ell(\mathbf{w}_t; \mathbf{x}, \mathbf{y})\|_2 \leq B$ for any t , $\mathbf{x}$, and $\mathbf{y}$
- $\bar{\mathbf{w}}_t = \frac{\mathbf{w}_1 + \dots + \mathbf{w}_t}{t}$
- The runtime is $O(1/\sqrt{t})$ if we choose $\eta = \frac{\|\mathbf{w}_0 - \mathbf{w}^*\|_2}{B\sqrt{T}}$, independent of the data set size n !

Mini-batch stochastic gradient descent

1. Sample a subset S_t from a data set S .
2. $\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t \nabla \frac{1}{|S_t|} \sum_{\mathbf{x}, \mathbf{y} \in S_t} \ell(\mathbf{w}_t; \mathbf{x}, \mathbf{y})$
 - The random sampling maintains $\mathbb{E}_{S_t \sim U(S)^t} \left[\nabla \frac{1}{|S_t|} \sum_{\mathbf{x}, \mathbf{y} \in S_t} \ell(\mathbf{w}; \mathbf{x}, \mathbf{y}) \right] = \nabla L(\mathbf{w})$.
3. Go to 1 until the solution is satisfactory.

Common practice and terminology

Common practice and terminology

- When doing SGD, permute the data and then go in serial.

Common practice and terminology

- When doing SGD, permute the data and then go in serial.
- A pass over the data set is called an **epoch**.

Common practice and terminology

- When doing SGD, permute the data and then go in serial.
- A pass over the data set is called an **epoch**.
- When doing mini-batch SGD, remember to normalise by the batch size.

Common practice and terminology

- When doing SGD, permute the data and then go in serial.
- A pass over the data set is called an **epoch**.
- When doing mini-batch SGD, remember to normalise by the batch size.
- With a larger batch size, we go over an epoch faster.

Common practice and terminology

- When doing SGD, permute the data and then go in serial.
- A pass over the data set is called an **epoch**.
- When doing mini-batch SGD, remember to normalise by the batch size.
- With a larger batch size, we go over an epoch faster.
- Use the largest batch size you can afford.