

Optimization Methods for Deep Learning



Gradient descent for non-convex optimization

Descent Lemma: Let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ be twice differentiable, and $\|\nabla^2 f\|_2 \leq \beta$. Then setting the learning rate $\eta = 1/\beta$, and applying gradient descent, $x_{t+1} = x_t - \eta \nabla f(x_t)$, we have:

$$f(x_t) - f(x_{t+1}) \geq \frac{1}{2\beta} \|\nabla f(x_t)\|_2^2.$$

Converging to stationary points

Theorem: In $T = O(\frac{\beta}{\epsilon^2})$ iterations, we have $\|\nabla f(x)\|_2 \leq \epsilon$.

Gradient Descent for Quadratic Functions

$$\text{optimal } x^* = 0 \quad / \quad \text{init: } x_0$$

Problem: $\min_x \frac{1}{2} x^T A x$ with $A \in \mathbb{R}^{d \times d}$ being positive-definite.

Theorem: Let $\lambda_{\max}$ and $\lambda_{\min}$ be the largest and the smallest eigenvalues of A . If we set $\eta \leq \frac{1}{\lambda_{\max}}$, we have

$$\|x_t\|_2 \leq (1 - \eta \lambda_{\min})^t \|x_0\|_2$$

$$\begin{aligned} \|x_{t+1}\|_2 &= \|x_t - \eta A x_t\|_2 \\ &= \|(I - \eta A) x_t\|_2 \\ &\leq \|I - \eta A\|_2 \cdot \|x_t\|_2 \\ &\leq (1 - \eta \lambda_{\min}) \cdot \|x_t\|_2 \\ &\leq (1 - \eta \lambda_{\min})^{t+1} \|x_0\|_2 \end{aligned}$$

$\circ$ linear convergence

If want $\|x_T\|_2 \leq \varepsilon$
we set $\eta = \frac{1}{\lambda_{\max}}$
 $\Rightarrow$ need $O\left(\frac{\lambda_{\max}}{\lambda_{\min}} \log\left(\frac{1}{\varepsilon}\right)\right)$
call $\kappa = \frac{\lambda_{\max}}{\lambda_{\min}}$
condition number

can be generalized to
strongly convex function

Momentum: Heavy-Ball Method (Polyak '64)

$$\begin{aligned} v_0 &= 0 \\ v_1 &= -\nabla f(x_0) \end{aligned}$$

Problem: $\min_x f(x)$

Method: $v_{t+1} = -\nabla f(x_t) + \beta v_t$

$$x_{t+1} = x_t + \eta v_{t+1}$$

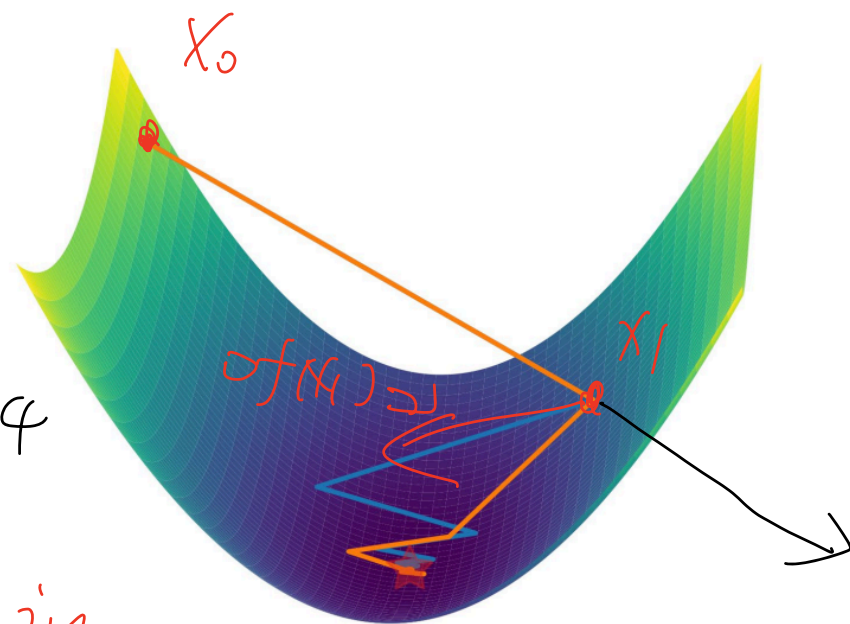
For quadratic optimization

$$\sqrt{K} \log\left(\frac{1}{\epsilon}\right)$$

VS. gradient $K \cdot \log\left(\frac{1}{\epsilon}\right)$

$K: 10^8 \Rightarrow \text{improv } 10^8 \rightarrow 10^4$

However, does not hold in
general strongly convex functions



Momentum: Nesterov Acceleration (Nesterov '89)

Problem: $\min_x f(x)$

look ahead

λ^-

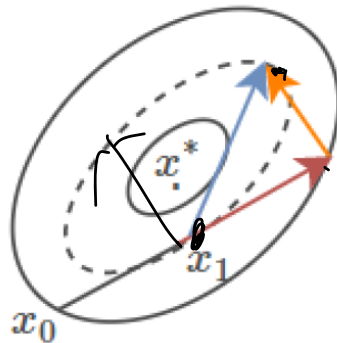
Method: $v_{t+1} = -\nabla f(x_t + \beta v_t) + \beta v_t$
 $x_{t+1} = x_t + \eta v_{t+1}$

For general strongly convex functions

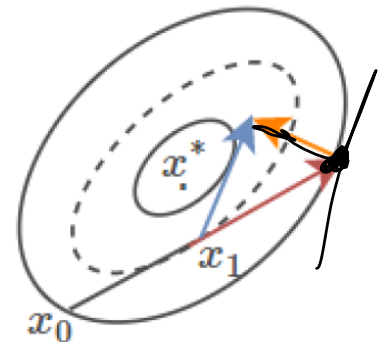
$\sqrt{n} \cdot \log\left(\frac{1}{\epsilon}\right)$

tight

Polyak's Momentum



Nesterov Momentum



$$x \in \mathbb{R}^d$$

$$O(d^2)$$

$$\text{inverse } O(d^3)$$

Newton's Method

$$d \times d$$

Newton's Method: $x_{t+1} = x_t - \eta (\underbrace{\nabla^2 f(x_t)}_O) \nabla f(x_t)$

GD: $x_{t+1} = x_t - \eta \nabla f(x_t)$

$$f(x+\Delta) \approx f(x) + \Delta^T \nabla f(x) + \frac{1}{2} \|\Delta\|_2^2$$

$$\Delta^* = -\nabla f(x)$$

Newton

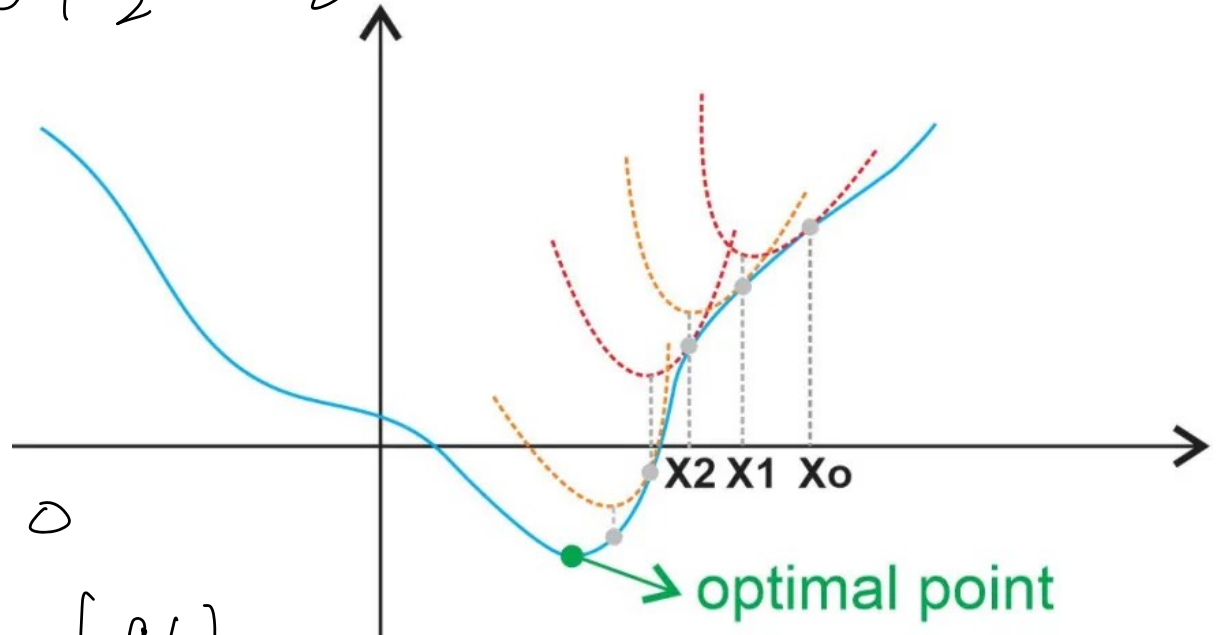
$$f(x+\Delta) \approx f(x) + \Delta^T \nabla f(x) + \frac{1}{2} \Delta^T \nabla^2 f(x) \Delta$$

$$\Rightarrow \Delta^* = -(\nabla^2 f(x))^{-1} \nabla f(x)$$

Theorem:

$$O(\log \log (\frac{1}{\epsilon}))$$

quadratic convergence



AdaGrad (Duchi et al. '11)

support diagonal

Newton Method: $x_{t+1} = x_t - \eta(\nabla^2 f(x_t))^{-1} \nabla f(x_t)$

AdaGrad: separate learning rate for every parameter

diagonal

$$x_{t+1} = x_t - \eta(\underbrace{G_{t+1} + \epsilon I}_{\text{invertible}})^{-1} \nabla f(x_t), (G_t)_{ii} = \sqrt{\sum_{j=1}^{t-1} (\nabla f(x_t)_i)^2}$$

in general
pre-conditioner

RMSProp (Hinton et al. '12)

AdaGrad: separate learning rate for every parameter

$$x_{t+1} = x_t - \eta(G_{t+1} + \epsilon I)^{-1} \nabla f(x_t), \quad \underbrace{(G_t)_{ii}}_{\text{strictly increasing}} = \sqrt{\sum_{j=1}^{t-1} (\nabla f(x_t)_i)^2}$$

RMSProp: exponential weighting of gradient norms

$$x_{t+1} = x_t - \eta(G_{t+1} + \epsilon I)^{-1/2} \nabla f(x_t),$$
$$\underbrace{(G_{t+1})_{ii}}_{\text{exponential weighting}} = \beta(G_t)_{ii} + (1 - \beta)(\nabla f(x_t)_i)^2$$

$0 < \beta < 1$

$\beta^{t+1} \nabla f(x_0)$

AdaDelta (Zeiler '12)

$$\frac{\partial f}{\partial x} : \frac{1}{\text{unit of } x}$$

RMSProp:

$$\left(\frac{1}{\text{unit of } x} \right)^2$$

unit of x

unit of loss

$$x_{t+1} = x_t - \eta (G_{t+1} + \epsilon I)^{-1/2} \nabla f(x_t),$$

$$(G_{t+1})_{ii} = \beta (G_t)_{ii} + (1 - \beta) (\nabla f(x_t)_i)^2$$

AdaDelta:

$\propto$ unit of x

$$x_{t+1} = x_t - \eta \Delta x_t,$$

$$\Delta x_t = \sqrt{u_t + \epsilon} \cdot (G_{t+1} + \epsilon I)^{-1/2} \nabla f(x_t)$$

$$(G_{t+1})_{ii} = \beta (G_t)_{ii} + (1 - \beta) (\nabla f(x_t)_i)^2,$$

$$u_{t+1} = \rho u_t + (1 - \rho) \|\Delta x_t\|_2^2$$

unitless

$$0 < \rho < 1$$

ρ close to

0.9, 0.99

GD: $\Delta x_t : \frac{1}{\text{unit of } x}$

Newton: $\Delta x_t \propto \left(\frac{\partial^2 f}{\partial x^2} \right)^{-1} \cdot \frac{\partial f}{\partial x} \propto \text{unit of } x$

Adam (Kingma & Ba '14)

Momentum:

$$v_{t+1} = -\nabla f(x_t) + \beta v_t, x_{t+1} = x_t + \eta v_{t+1}$$

RMSProp: exponential weighting of gradient norms

$$x_{t+1} = x_t - \eta (G_{t+1} + \epsilon I)^{-1} \nabla f(x_t),$$
$$(G_t)_{ii} = \beta (G_t)_{ii} + (1 - \beta) (\nabla f(x_t)_i)^2$$

Adam

$$v_{t+1} = \beta_1 v_t + (1 - \beta_1) \nabla f(x_t)$$
$$(G_{t+1})_{ii} = \beta_2 (G_t)_{ii} + (1 - \beta_2) (\nabla f(x_t)_i)^2$$
$$x_{t+1} = x_t - \eta (G_{t+1} + \epsilon I)^{-1/2} v_{t+1}$$

Default choice nowadays.

Important Techniques in Neural Network Training



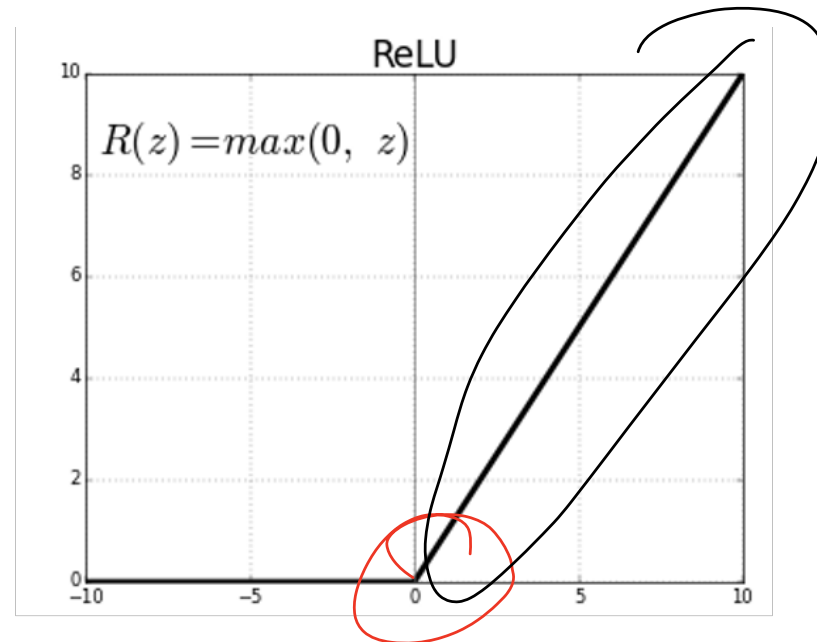
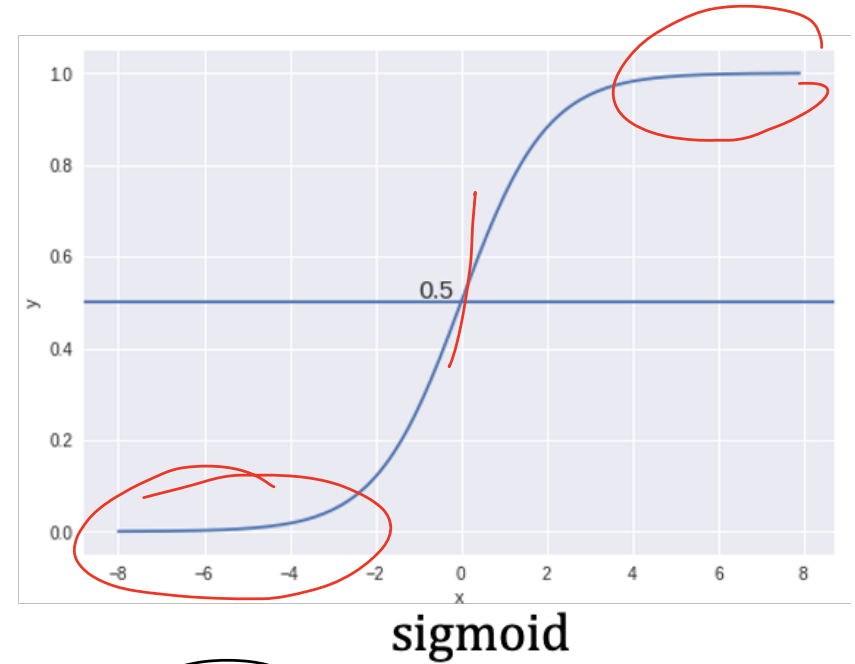
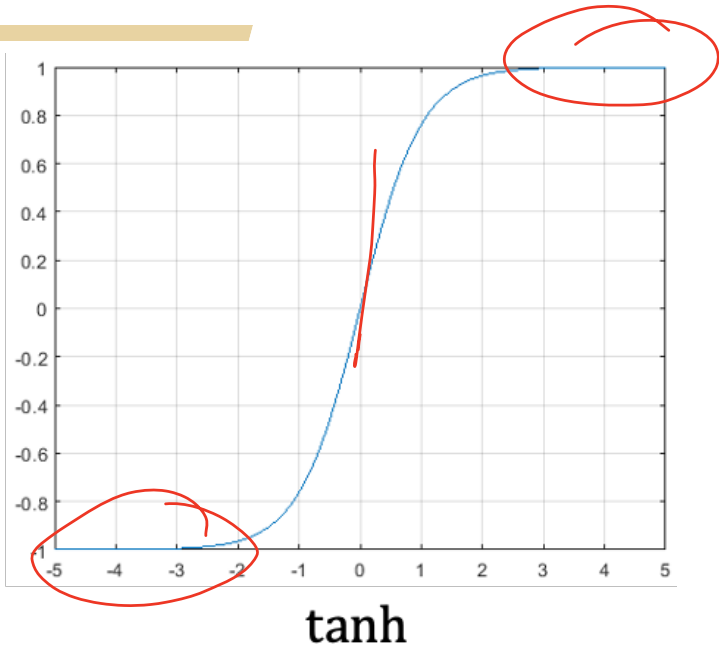
- $$f(x, w_1, \dots, w_{H+1}) = w_{H+1} \sigma(w_H \dots \sigma(w_1 x) \dots)$$
- $$\frac{\partial f}{\partial w_n} = (w_{H+1} A_H \dots w_{n+1} A_n)^T (A_{n-1} w_{n-1} \dots w_1 x)$$

$$A_h = \text{diag} (\sigma^1(w_h) \sigma^1(w_1 x) \dots)$$

If W $\left\{ \begin{array}{l} \text{magnitude large} \\ \text{magnitude small} \end{array} \right. \rightarrow \begin{array}{l} \text{gradient exploding} \\ \text{gradient vanishing} \end{array}$

A:

Activation Functions



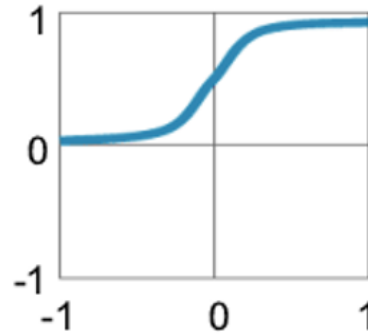
Rectified Linear Unit

= /

Activation Function

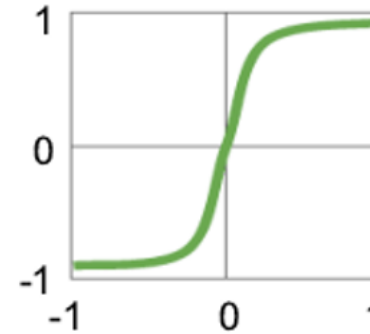
Traditional Non-Linear Activation Functions

Sigmoid



$$y = 1 / (1 + e^{-x})$$

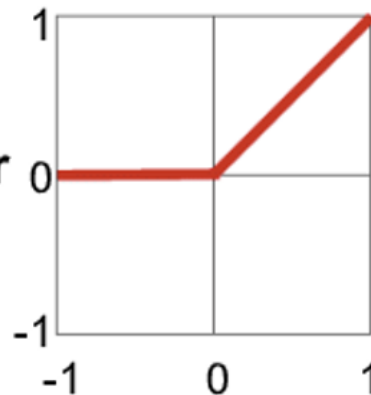
Hyperbolic Tangent



$$y = (e^x - e^{-x}) / (e^x + e^{-x})$$

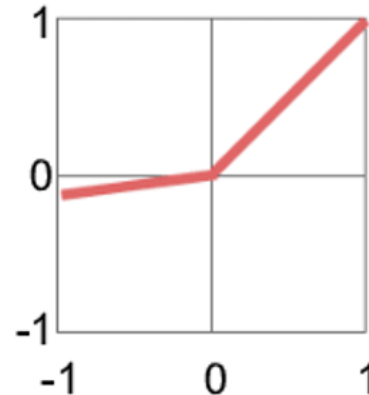
Modern Non-Linear Activation Functions

Rectified Linear Unit (ReLU)



$$y = \max(0, x)$$

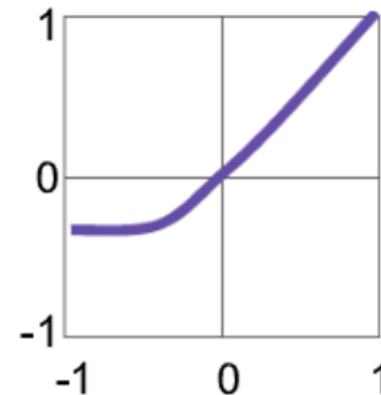
Leaky ReLU



$$y = \max(\alpha x, x)$$

α = small const. (e.g. 0.1)

Exponential LU



$$y = \begin{cases} x, & x \geq 0 \\ \alpha(e^x - 1), & x < 0 \end{cases}$$

Initialization

W_0

$$(W^T x - y)^2$$

- Zero-initialization
- Large initialization
- Small initialization

$W_{ij} \sim \text{Gaussian}$

Random

$\sim \text{Unit}$

$$E[W_{ij}] = 0$$

- Design principles:
 - Zero activation mean

- Activation variance remains same across layers

Gaussian : $N(0, \sigma^2)$, variance = 1

Unit : Unit $[-a, a]$, variance = 1