

Optimization Methods for Deep Learning



Gradient descent for non-convex optimization

$$\|A\|_2 = \max_{\|v\|_2=1} v^T A v, \quad v \in \mathbb{R}^d$$

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 $\eta = \underline{1/\beta}$, and applying gradient descent, $x_{t+1} = x_t - \eta \nabla f(x_t)$, we have: can use $\frac{1}{\beta}$

$$f(x_t) - f(x_{t+1}) \geq \frac{1}{2\beta} \|\nabla f(x_t)\|_2^2. > 0$$

pf: by Taylor expansion & mean-value theorem $\Rightarrow$

$$f(x+\Delta) = f(x) + \Delta^T \nabla f(x) + \frac{1}{2} \Delta^T \nabla^2 f(y) \Delta \quad \text{for some } y$$

$$\Delta^T \nabla^2 f(y) \Delta \leq \|\nabla^2 f(y)\|_2 \cdot \|\Delta\|_2^2 \leq \beta \cdot \|\Delta\|_2^2$$

choose $\Delta = -\eta \nabla f(x_t)$

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

$$= f(x_t) - \frac{1}{2\beta} \|\nabla f(x_t)\|_2^2$$

Converging to stationary points

$$\eta = \frac{1}{\beta}$$

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

$$\text{Pf: } f(x_{t+1}) \leq f(x_t) - \frac{\eta}{2} \|\nabla f(x_t)\|_2^2$$

$$\text{sum over } t=0, \dots, T-1$$
$$\sum_{t=1}^T f(x_t) \leq \sum_{t=0}^{T-1} f(x_t) - \frac{\eta}{2} \sum_{t=0}^{T-1} \|\nabla f(x_t)\|_2^2$$

$$\Rightarrow f(x_T) \leq f(x_0) - \frac{\eta}{2} \sum_{t=0}^{T-1} \|\nabla f(x_t)\|_2^2$$

$$\Rightarrow \sum_{t=0}^{T-1} \|\nabla f(x_t)\|_2^2 \leq \frac{2(f(x_0) - f(x_T))}{\eta}$$

$$T, \min_{0 \leq t \leq T-1} \|\nabla f(x_t)\|_2 \leq$$

$$\min_{0 \leq t \leq T-1} \|\nabla f(x_t)\|_2 \leq \sqrt{\frac{2(f(x_0) - f(x_T))}{\eta \cdot T}} \leq \epsilon$$

$$T = O\left(\frac{\beta}{\epsilon^2}\right)$$

Gradient Descent for Quadratic Functions

$$x^* = 0$$

$$\nabla f(x) = Ax$$

$$\lambda_{\min} > 0$$

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

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

$$\begin{aligned} \|x_{t+1}\|_2 &= \|x_t - \eta A x_t\|_2 \\ &= \|(I - \eta A) x_t\|_2 \end{aligned}$$

$$\begin{aligned} \lambda_{\max}(\eta A) &\leq 1 \quad \Rightarrow \|I - \eta A\|_2 \cdot \|x_t\|_2 \\ &\leq (1 - \eta \lambda_{\min}) \cdot \|x_t\|_2 \\ &\leq \underbrace{(1 - \eta \lambda_{\min})^{t+1}}_{\text{linear convergence}} \|x_0\|_2 \end{aligned}$$

Implication:

$$\text{If } \|x_T\|_2 \leq \epsilon$$

$$\text{if } \eta = \frac{1}{\lambda_{\max}}$$

$$\Rightarrow T = O\left(\frac{\lambda_{\max}}{\lambda_{\min}} \cdot \log\left(\frac{1}{\epsilon}\right)\right)$$

$$\kappa = \frac{\lambda_{\max}}{\lambda_{\min}} \quad \text{condition number}$$

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

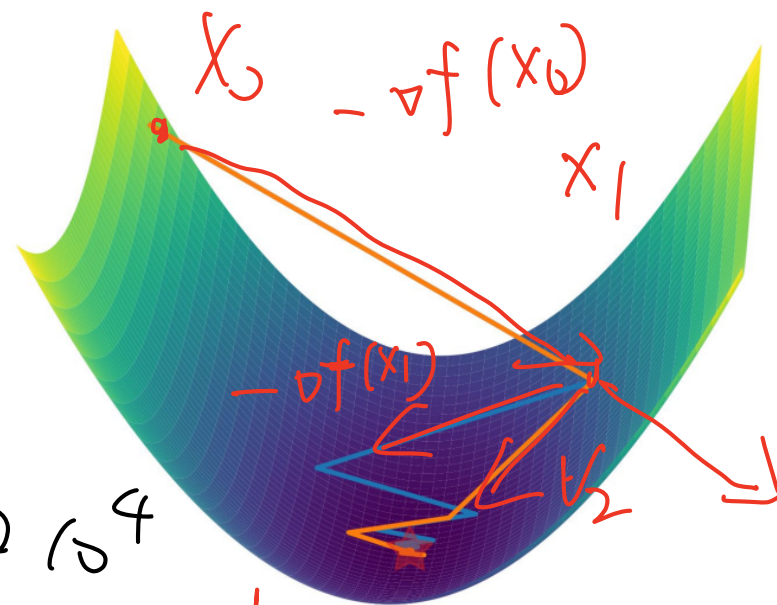
$$\mathcal{O}\left(\sqrt{K} \log\left(\frac{1}{\epsilon}\right)\right)$$

V.S. Gradient descent

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

$$K = 10^8 \Rightarrow \text{improve } 10^8 \rightarrow 10^4$$

However, does not hold for general strongly convex function



Momentum: Nesterov Acceleration (Nesterov '89)

Problem: $\min_x f(x)$

$$v_t, -\nabla f(x_t)$$

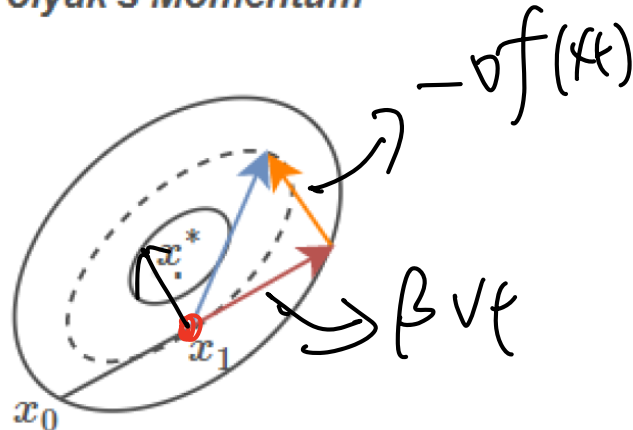
Method: $v_{t+1} = -\nabla f(x_t + \beta v_t) + \beta v_t$

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

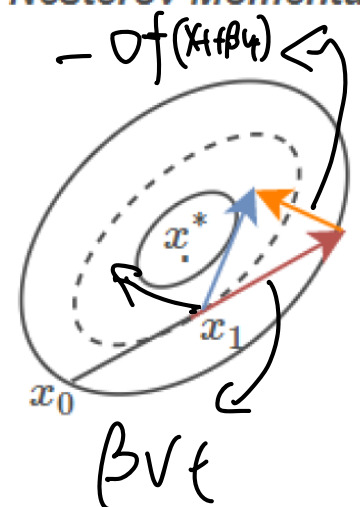
For general strongly convex function
 $\sqrt{K} \cdot \log\left(\frac{1}{\epsilon}\right)$

tight

Polyak's Momentum



Nesterov Momentum



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

$$O(d^2) \text{ inverse } O(d^3)$$

Newton's Method

$$dx dx$$

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

• Q1): $x_{t+1} = x_t - \eta \nabla f(x_t)$

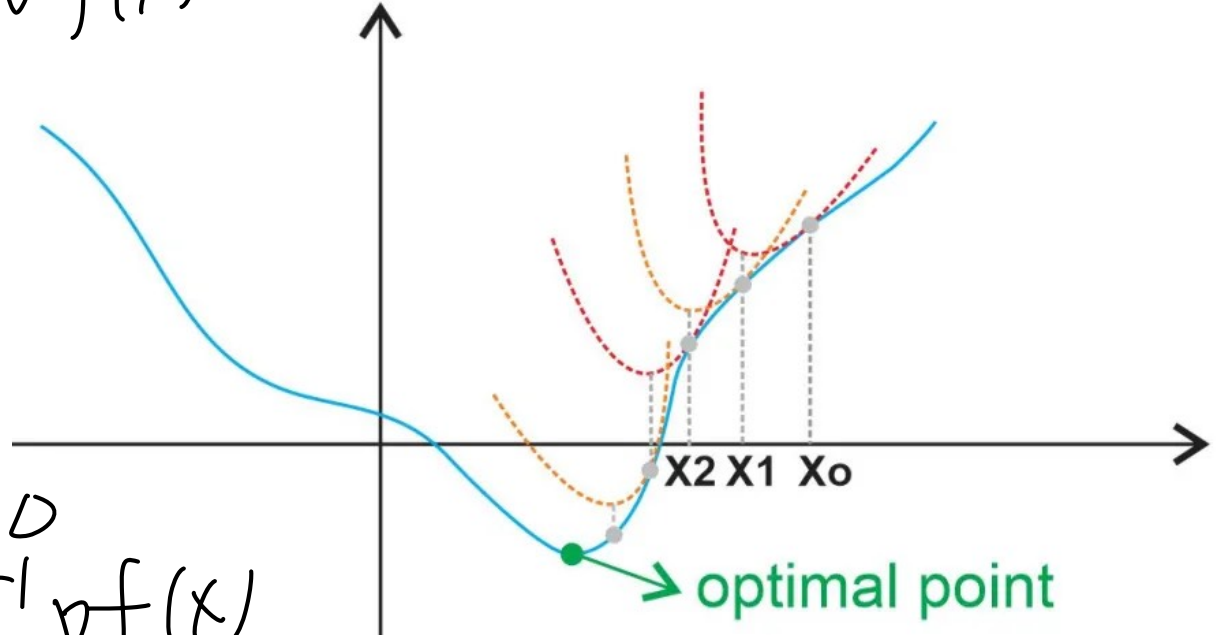
$$f(x+\Delta) \approx f(x) + \Delta^T \nabla f(x) + \frac{1}{2} \Delta^T \nabla^2 f(x) \Delta$$

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

• 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: $O(\log \log(\frac{1}{\epsilon}))$ vs. $\log(\frac{1}{\epsilon})$
 qua diversi convergenze

AdaGrad (Duchi et al. '11)

suppose diagonal
 $\mathcal{O}(d)$

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

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

make invertible

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

preconditioning

RMSProp (Hinton et al. '12)

AdaGrad: separate learning rate for every parameter

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

Strictly increasing

RMSProp: exponential weighting of gradient norms

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

$0 < \beta < 1$, close to 1
exponential weighting
EMA

AdaDelta (Zeiler '12)

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

unit of x

RMSProp: $\frac{1}{(\text{unit of } x)^2}$

$\Rightarrow$ unit less

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

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

AdaDelta:

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

unit less / unit of x

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

unit of x

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

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

$0 < \rho < 1$ unit of x^2

$\nabla^2 f : - \text{of } (x^2) : \frac{1}{\text{unit of } x}$

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

Adam (Kingma & Ba '14)

Momentum: heavy ball

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

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



Gradient Explosion / Vanishing

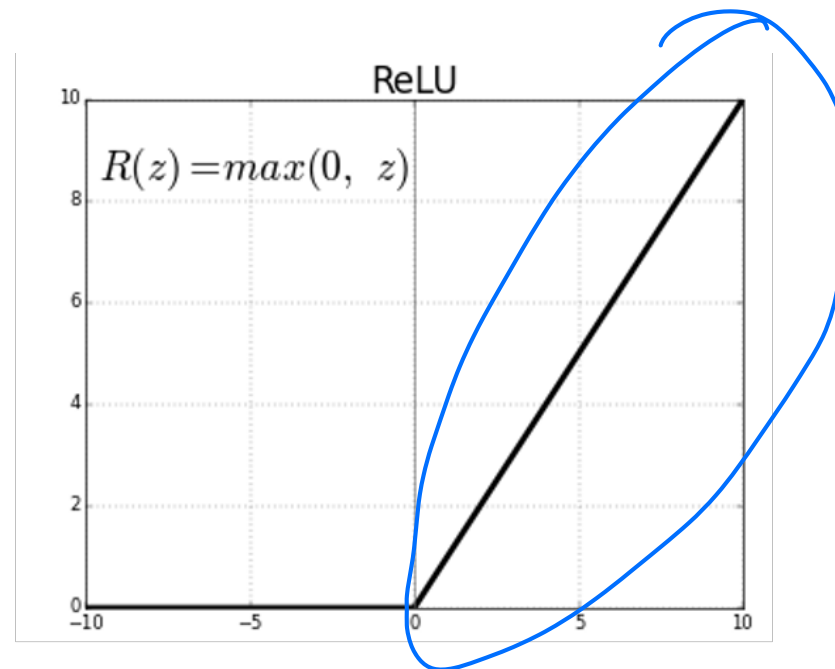
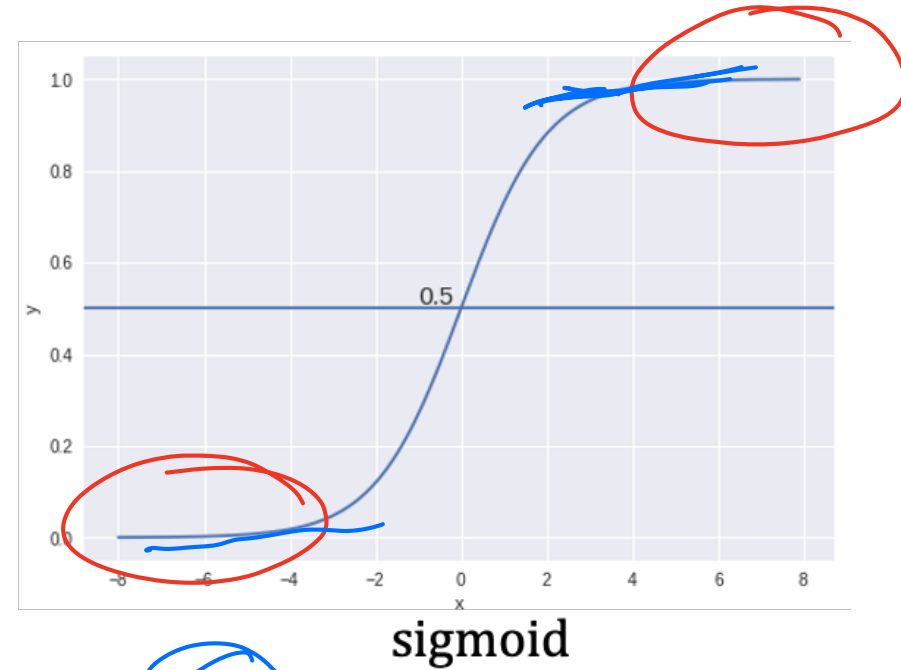
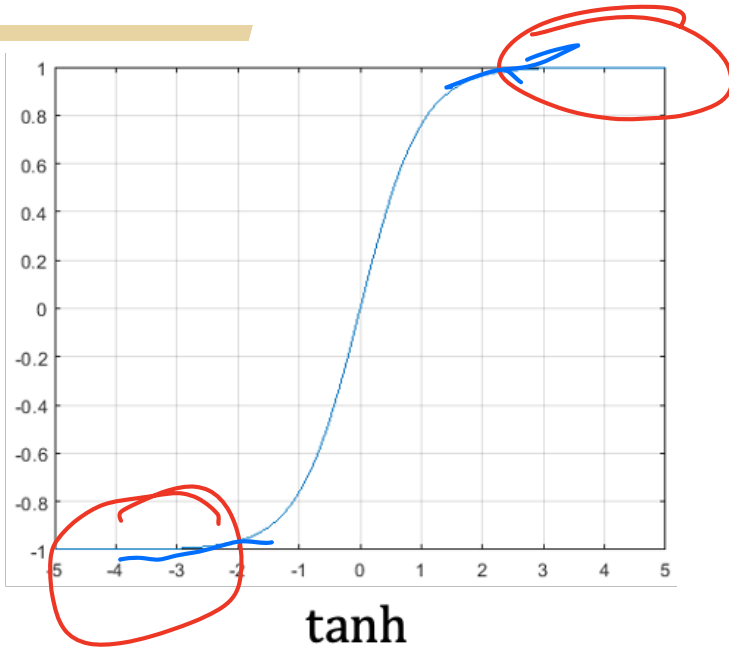
- Deeper networks are harder to train:
 - Intuition: gradients are products over layers
 - Hard to control the learning rate

$$f(x, w_1, \dots, w_{H+1}) = w_{H+1} \sigma(w_H \dots \sigma(w_1 x) \dots)$$
$$\frac{\partial f}{\partial w_h} = (w_{H+1} A_{H+1} \dots w_{h+1} A_h)^T (A_{h-1} w_{h-1} \dots w_1 x)$$
$$A_h = \text{diag}(\sigma'(w_h \sigma(\dots \sigma(w_1 x) \dots))$$

If $w_1, \dots, w_{H+1}$ magnitude large $\rightarrow$ gradient explosion
 $w_1, \dots, w_{H+1}$ magnitude small $\rightarrow$ gradient vanishing

Activation Functions

gradient vanishing



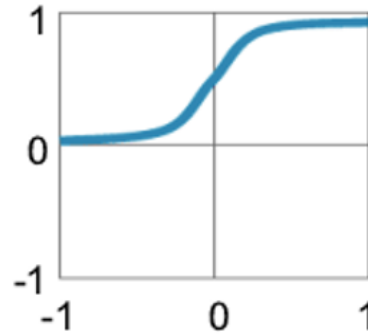
Rectified Linear Unit

$$G'(\cdot) = 1$$

Activation Function

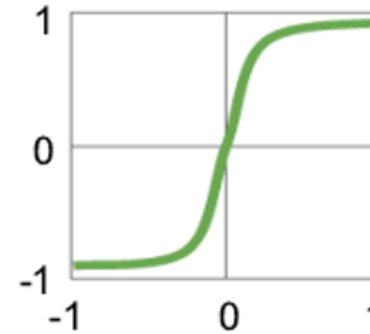
Traditional Non-Linear Activation Functions

Sigmoid



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

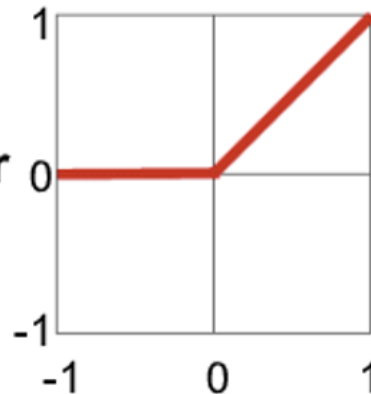
Hyperbolic Tangent



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

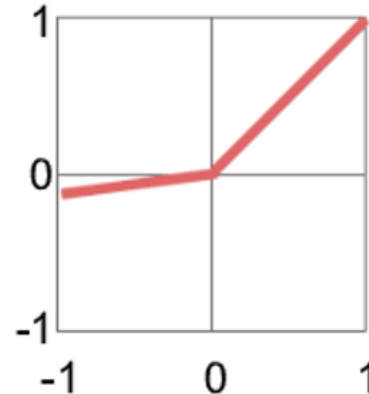
Modern Non-Linear Activation Functions

Rectified Linear Unit (ReLU)



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

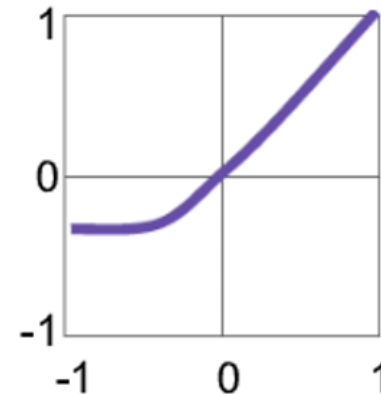
Leaky ReLU



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

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

Set $w_1 \dots, w_{H+1} = 0$ ~~$\times$~~ $\text{grad} = 0$

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

- Zero-initialization
- Large initialization $\Rightarrow$ all w large $\Rightarrow$ grad exploding
- Small initialization $\Rightarrow$ all w small $\Rightarrow$ grad vanishing

- Design principles:

- Zero activation mean

No prior

- Activation variance remains same across layers

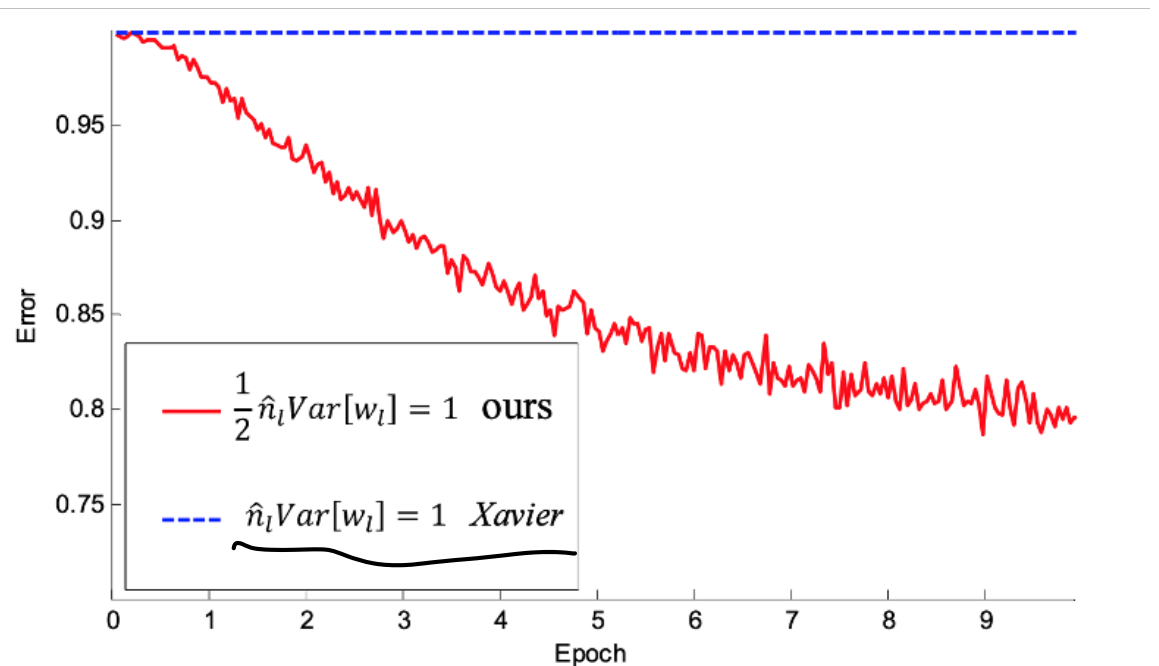
Keep every layer balanced

Kaiming Initialization (He et al. '15)

- $W_{ij}^{(h)} \sim \mathcal{N}\left(0, \frac{2}{d_h}\right)$.
- $b^{(h)} = 0$
- Designed for ReLU activation
- 30-layer neural network

$$W^{(h)} \in \mathbb{R}^{d_{h+1} \times d_h}$$

$$\mathcal{N}\left(0, \frac{1}{d_h}\right)$$



Kaiming Initialization (He et al. '15)

For layer h :

$$Z^h = W^h X^h$$

pre-activation

$$X^{h+1} = \sigma(Z^h)$$

$$Z_i^h = \sum_{j=1}^{d_h} W_{ij}^h \cdot X_j^{h-1}$$

Goal: Z^h : mean zero, same variance for all layers

$$\begin{aligned} \text{if } \mathbb{E}[W_{ij}^h] &= 0 \Rightarrow \mathbb{E}[Z_i^h] = \mathbb{E}\left[\sum_{j=1}^{d_h} W_{ij}^h X_j^{h-1}\right] \\ &= \sum_{j=1}^{d_h} \mathbb{E}[W_{ij}^h] \cdot X_j^{h-1} \\ &= 0 \end{aligned}$$

Kaiming Initialization (He et al. '15)

$$\begin{aligned}\text{Var}(z_i^h) &= d_h \text{Var}(w_{ij}^h \cdot x_j^h) \\ &= d_h \left[\underbrace{\text{Var}(w_{ij}^h) \cdot \text{Var}(x_j^h)}_{\text{red underline}} + \underbrace{\left(\mathbb{E}[w_{ij}^h]\right)^2 \cdot \text{Var}(x_j^h)}_{\text{red underline}} + \underbrace{\text{Var}(w_{ij}^h) \cdot \left(\mathbb{E}[x_j^h]\right)^2}_{\text{red underline}} \right]\end{aligned}$$

$$= d_h \text{Var}(w_{ij}^h) \cdot \mathbb{E}[(x_j^h)^2]$$

$$\begin{aligned}\mathbb{E}[(x_j^h)^2] &= \int_{-\infty}^{\infty} (x_j^h)^2 p(x_j^h) dx_j^h \\ &= \int_{-\infty}^{\infty} \max(0, z_j^{h-1})^2 p(z_j^{h-1}) dz_j^{h-1} \\ &= \int_0^{\infty} (z_j^{h-1})^2 p(z_j^{h-1}) dz_j^{h-1}\end{aligned}$$

(ReLU)

Kaiming Initialization (He et al. '15)

(symmetry of init around 0) $= \frac{1}{2} \int_{-\infty}^{\infty} (z_j^{h-1})^2 p(z_j^{h-1}) dz_j^{h-1}$

$$= \frac{1}{2} \text{Var}(z_j^{h-1})$$

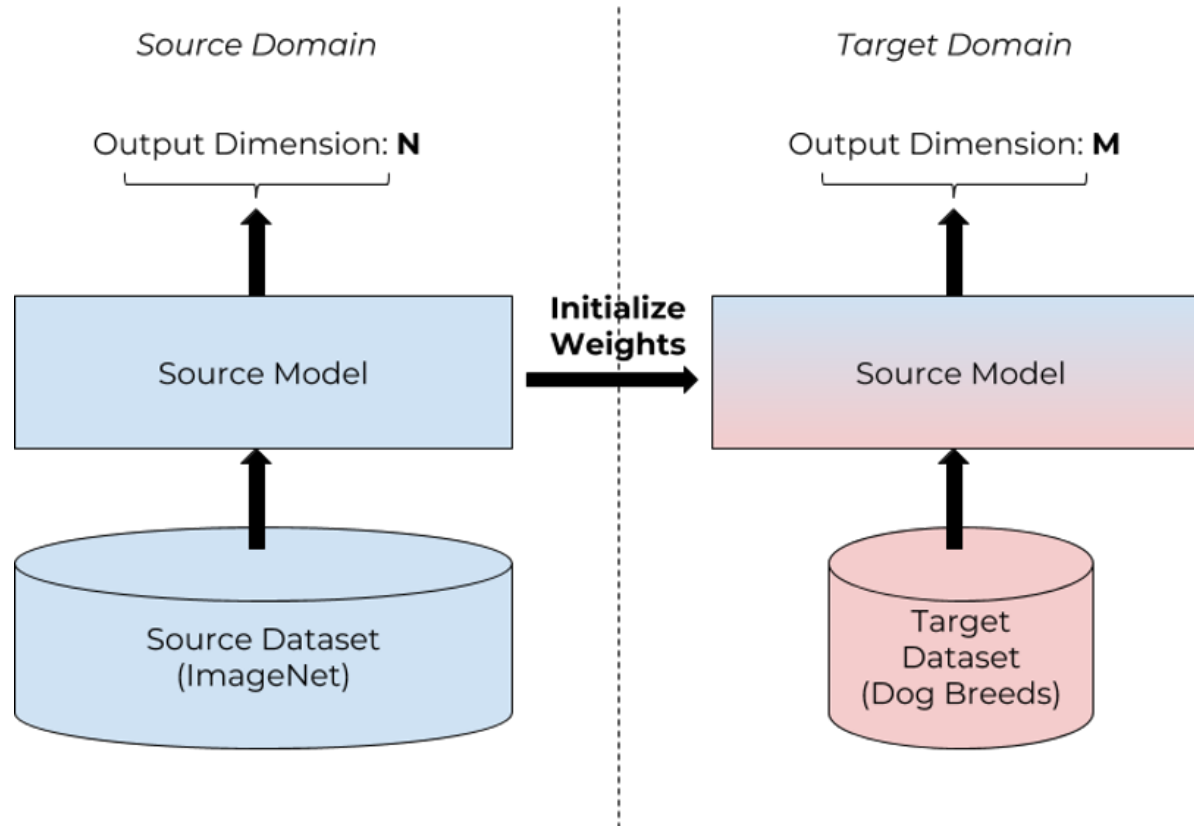
$$\text{Var}(z_j^h) = d_h \text{Var}(w_{ij}^h) \cdot \frac{1}{2} \text{Var}(z_j^{h-1})$$

Want $\text{Var}(z_j^h) = \text{Var}(z_j^{h-1})$

need: $d_h \cdot \text{Var}(w_{ij}^h) \cdot \frac{1}{2} = 1$
 $\text{Var}(w_{ij}^h) = \frac{2}{d_h}$

Initialization by Pre-training

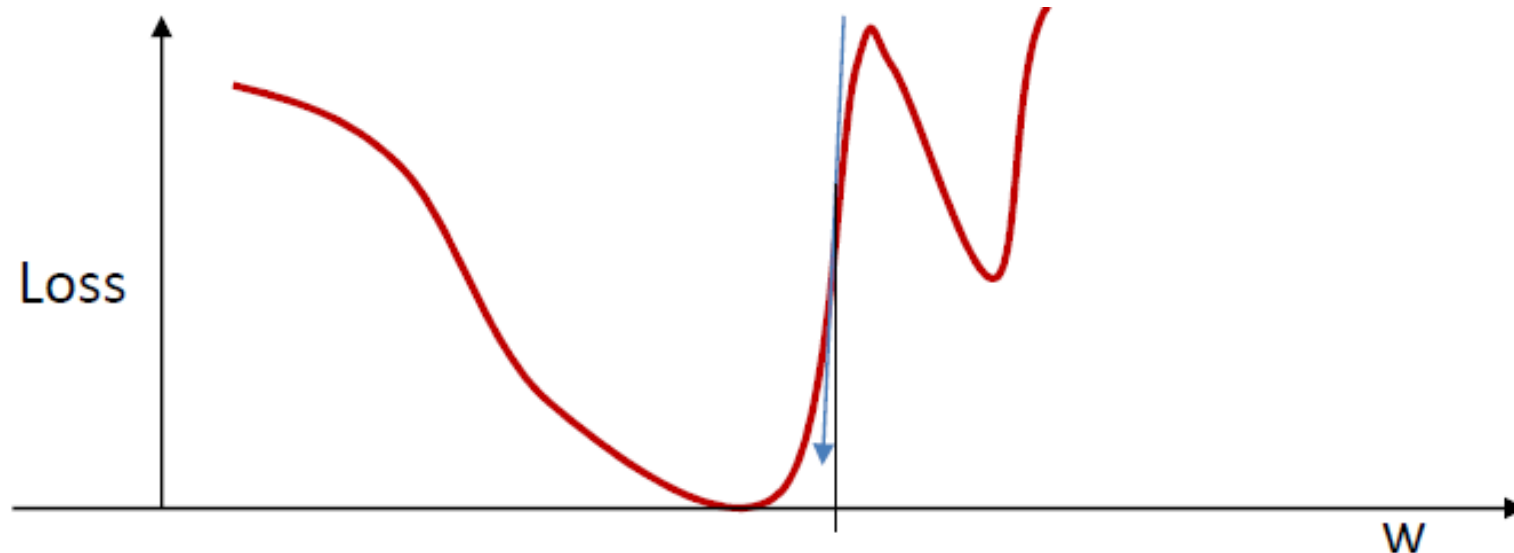
- Use a pre-trained network as initialization
- And then fine-tuning



Gradient Clipping

$$\text{if } \|\nabla f(x_t)\|_2 \geq \text{threshold}$$
$$g_t \leftarrow \frac{\nabla f(x_t)}{\|\nabla f(x_t)\|_2} \cdot \text{threshold}$$
$$x_{t+1} \leftarrow x_t - \eta g_t$$

- The loss can occasionally lead to a steep descent
- This result in immediate instability
- If gradient norm bigger than a threshold, set the gradient to the threshold.



Batch Normalization (Ioffe & Szegedy, '14)

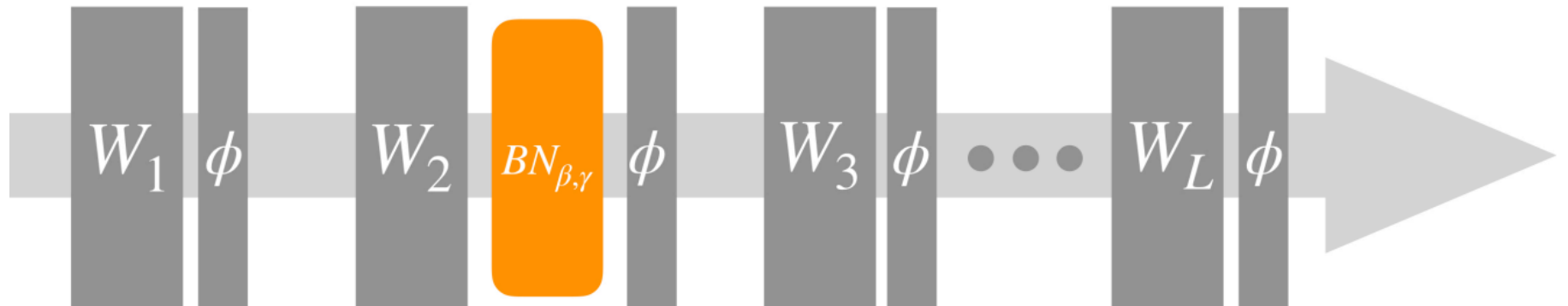
- **Normalizing/whitening** (mean = 0, variance = 1) the inputs is generally useful in machine learning.
 - Could normalization be useful at the level of hidden layers?
 - **Internal covariate shift**: the calculations of the neural networks change the distribution in hidden layers even if the inputs are normalized
- **Batch normalization** is an attempt to do that:
 - Each unit's **pre-activation** is normalized (mean subtraction, std division)
 - During training, mean and std is computed for each minibatch (can be backproped!)

Batch Normalization (Ioffe & Szegedy, '14)

Standard Network



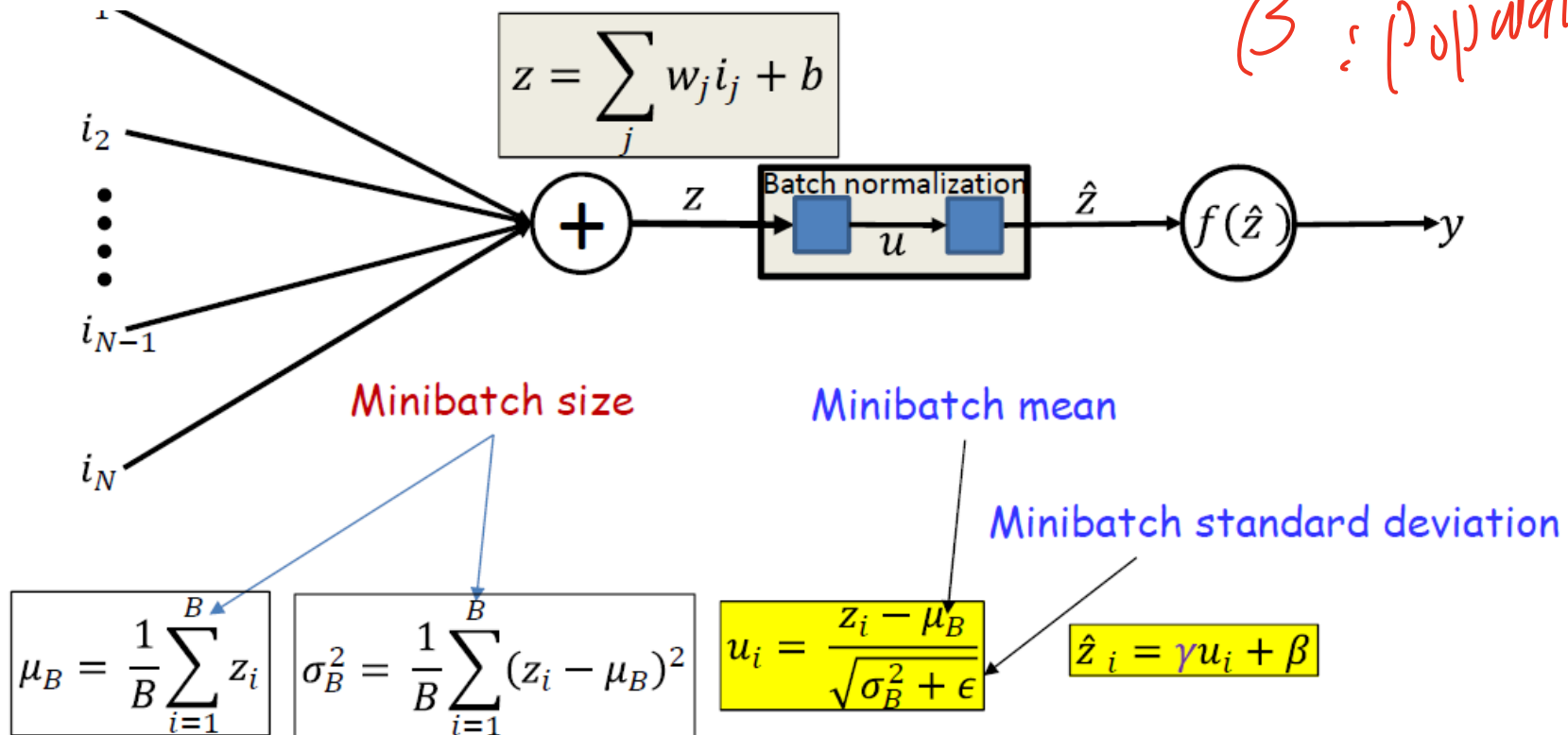
Adding a BatchNorm layer (between weights and activation function)



Batch Normalization (Ioffe & Szegedy, '14)

input $z_1, \dots, z_B$

γ : population std.
 β : population mean

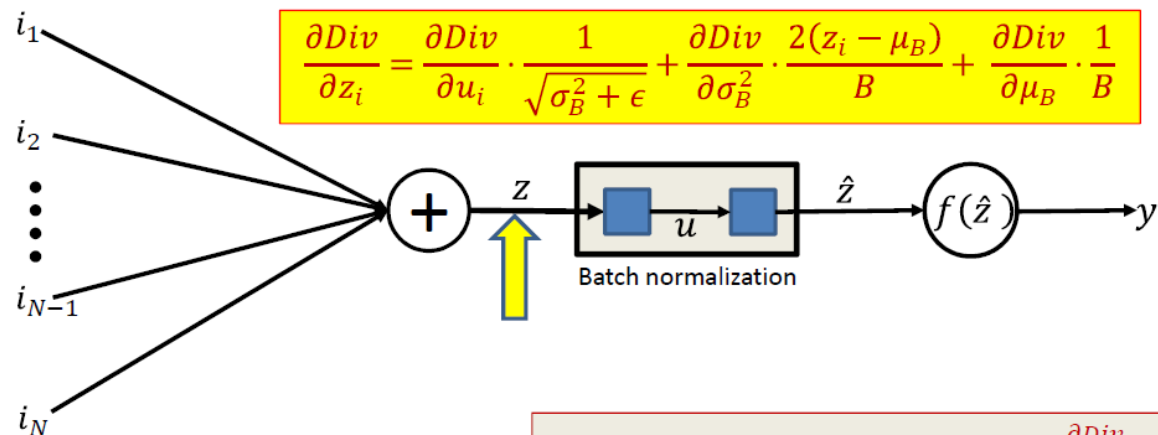


Batch Normalization (Ioffe & Szegedy, '14)

- BatchNorm at training time
 - Standard backprop performed for each single training data
 - Now backprop is performed over entire batch.

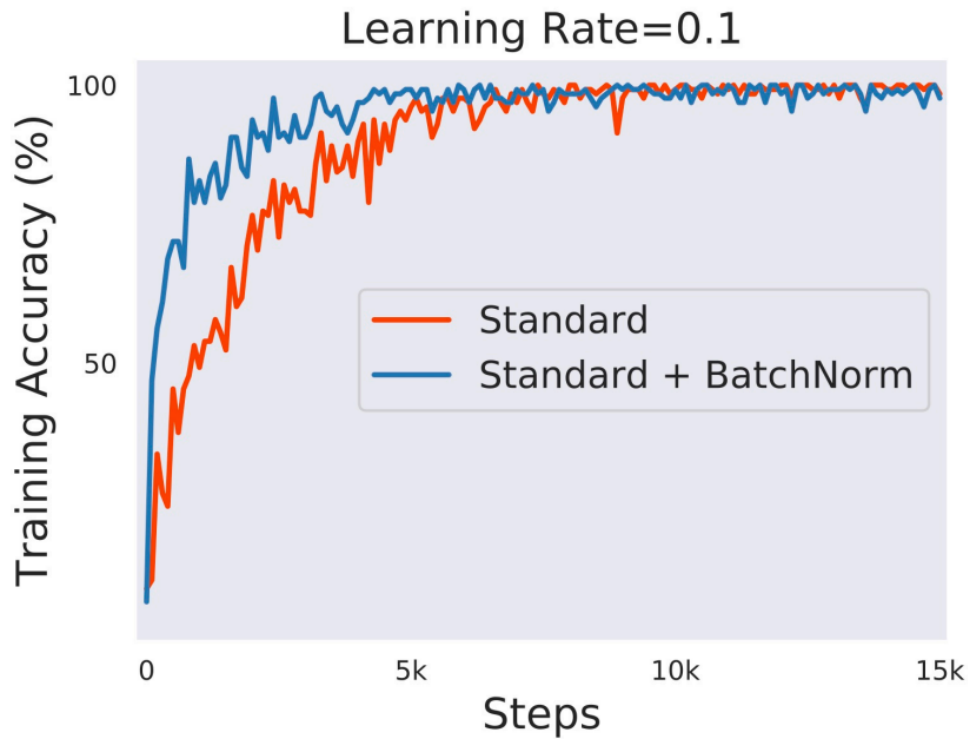
$$\frac{\partial \text{Div}}{\partial \sigma_B^2} = \frac{-1}{2} (\sigma_B^2 + \epsilon)^{-3/2} \sum_{i=1}^B \frac{\partial \text{Div}}{\partial u_i}$$

$$\frac{\partial \text{Div}}{\partial \mu_B} = \frac{-1}{\sqrt{\sigma_B^2 + \epsilon}} \sum_{i=1}^B \frac{\partial \text{Div}}{\partial u_i}$$



The rest of backprop continues from $\frac{\partial \text{Div}}{\partial z_i}$

Batch Normalization (Ioffe & Szegedy, '14)



What is BatchNorm actually doing?

- May not due to covariate shift (Santurkar et al. '18):
 - Inject non-zero mean, non-standard covariance Gaussian noise after BN layer: removes the whitening effect
 - Still performs well.
- Only training β, γ with random convolution kernels gives non-trivial performance (Frankle et al. '20)
- BN can use exponentially increasing learning rate! (Li & Arora '19)

More normalizations

- Layer normalization (Ba, Kiros, Hinton, '16)
 - Batch-independent
 - Suitable for RNN, MLP
- Weight normalization (Salimans, Kingma, '16)
 - Suitable for meta-learning (higher order gradients are needed)
- Instance normalization (Ulyanov, Vedaldi, Lempitsky, '16)
 - Batch-independent, suitable for generation tasks
- Group normalization (Wu & He, '18)
 - Batch-independent, improve BatchNorm for small batch size

Non-convex Optimization Landscape

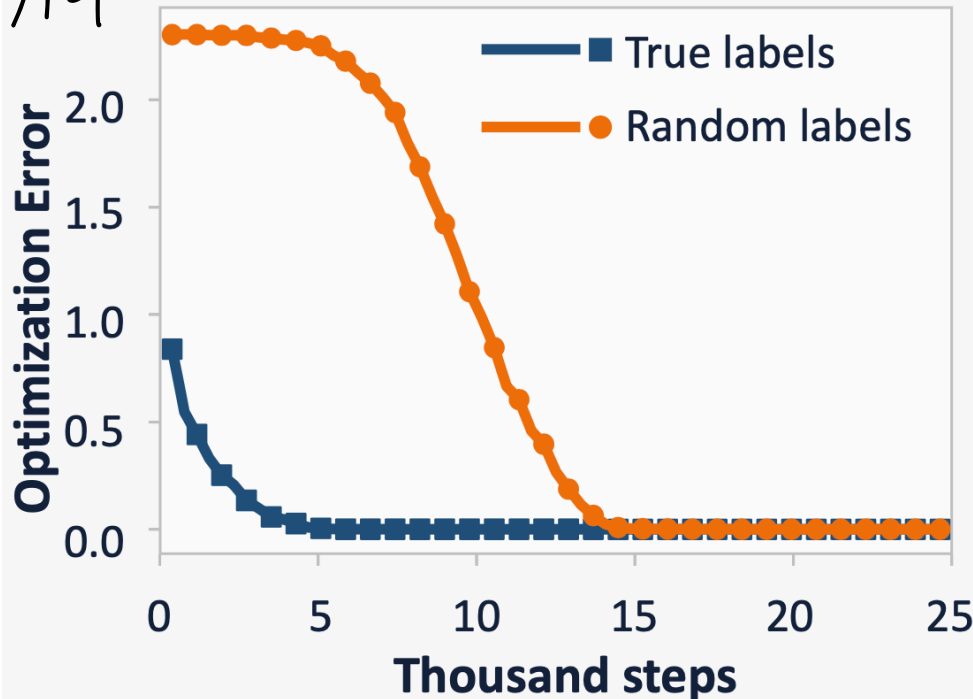
W

Gradient descent finds global minima

large NN
over-parameterized
of params

Practice: gradient descent

$$\theta(t+1) \leftarrow \theta(t) - \eta \frac{\partial L(\theta(t))}{\partial \theta(t)}$$



Optimization error $\rightarrow 0$ for both **true labels** and **random labels** !

$\Rightarrow$ $\exists$ NN
 $\Rightarrow$ 0 error

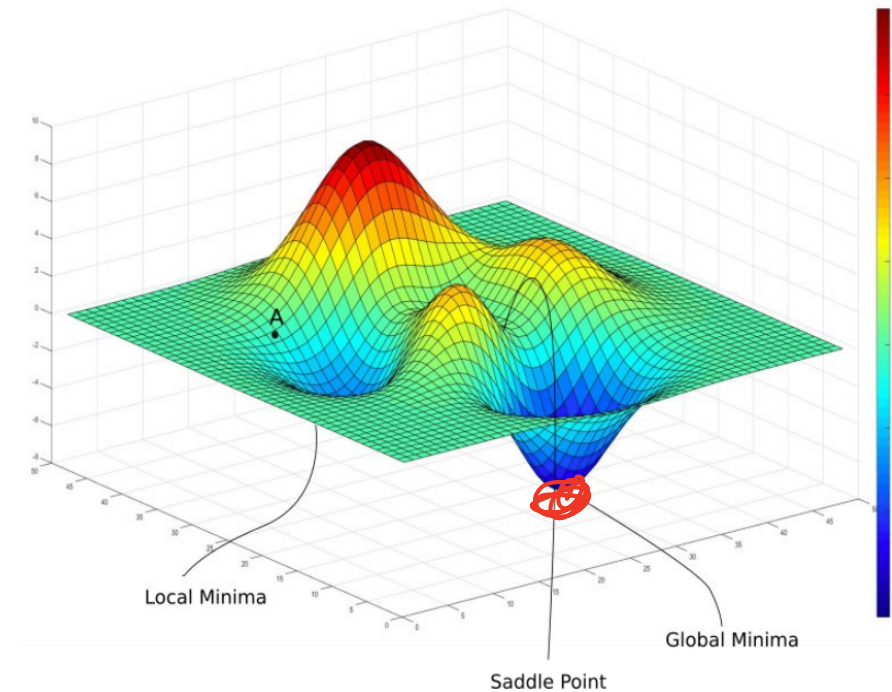
Zhang Bengio Hardt Recht Vinyals 2017

Understanding DL Requires Rethinking Generalization

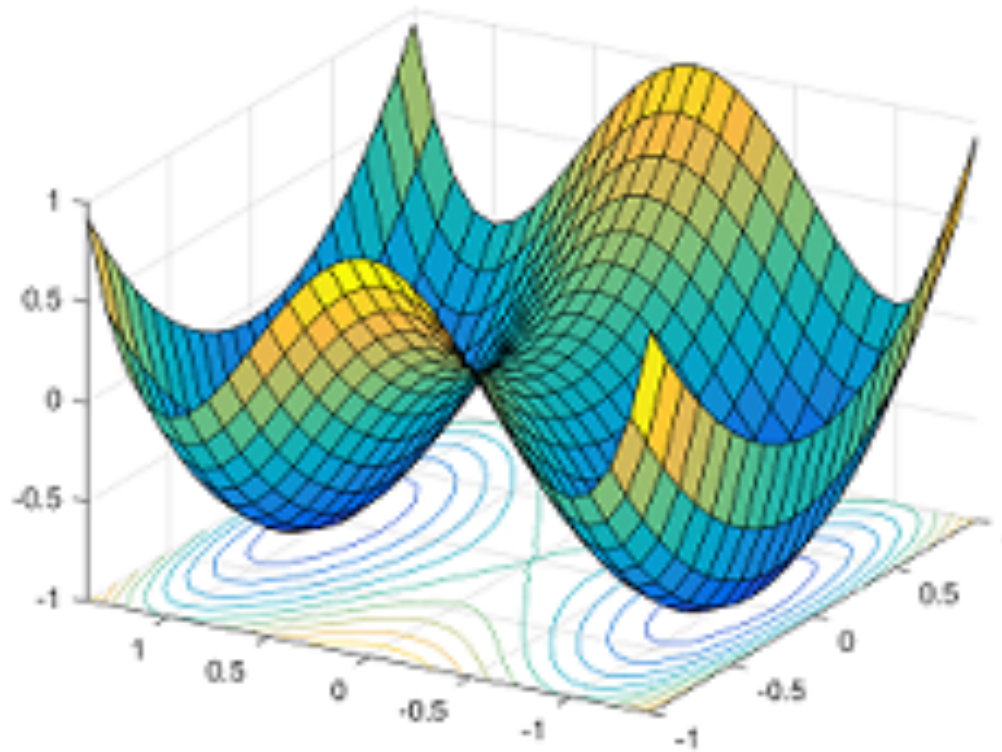
$\{(x_i, y_i)\}_{i=1}^n$
 $\{(x_i, \tilde{y}_i)\}_{i=1}^n$
 $\tilde{y}_i$: random numbers

Types of stationary points

- Stationary points: $x : \nabla f(x) = 0$
- Global minimum:
 $x : f(x) \leq f(x') \forall x' \in \mathbb{R}^d$
- Local minimum:
 $x : f(x) \leq f(x') \forall x' : \|x - x'\| \leq \epsilon$
- Local maximum:
 $x : f(x) \geq f(x') \forall x' : \|x - x'\| \leq \epsilon$
- Saddle points: stationary points that are not a local min/max

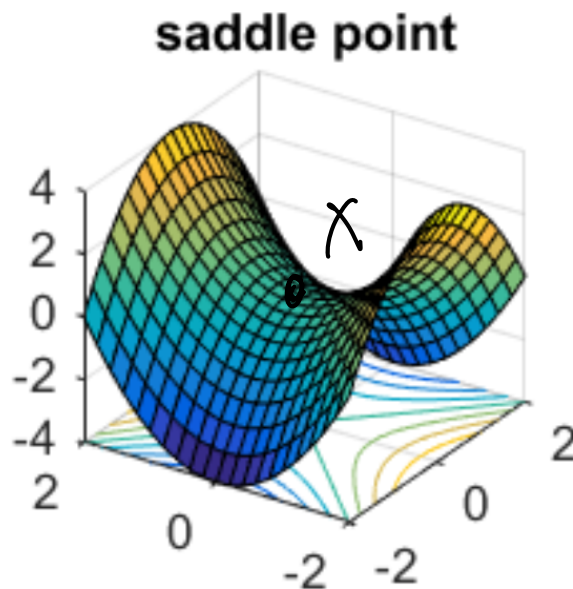


Landscape Analysis



- All local minima are global!
- Gradient descent can escape saddle points.

Strict Saddle Points (Ge et al. '15, Sun et al. '15)



local min: x

$\nabla^2 f(x)$ is
positive def

- Strict saddle point: a saddle point and $\lambda_{\min}(\underbrace{\nabla^2 f(x)}_A) < 0$

$$\text{min } \frac{1}{2} x^T A^T x, \quad \nabla f(x) = Ax$$

$$\begin{aligned} |v^T x_{t+1}| &= v^T (x_t - \eta_t A x_t) \\ &= v^T x_t - \eta_t \lambda_{\min} v^T x_t \\ &= \underbrace{(1 - \eta_t \lambda_{\min})}_{> 1} |v^T x_t| \end{aligned}$$

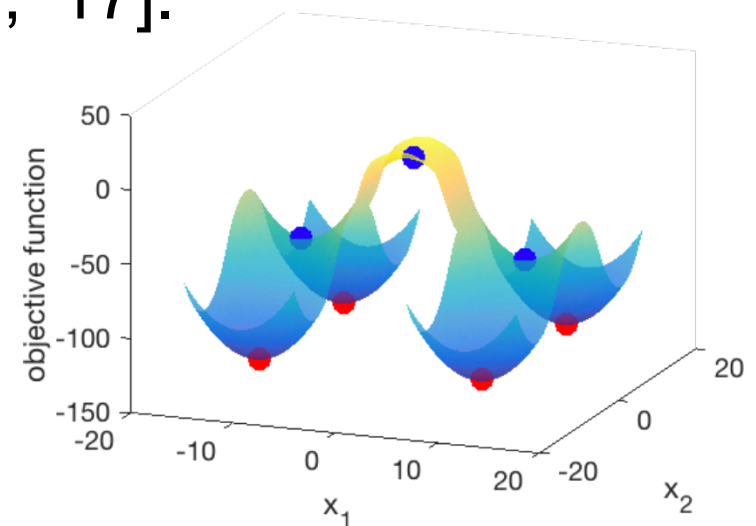
$x=0$ stationary point

v : eigen vector
(corresponds to
smallest eigenvalue)

Escaping Strict Saddle Points

- **Noise-injected** gradient descent can escape strict saddle points in polynomial time [Ge et al., '15, Jin et al., '17].
- Randomly initialized gradient descent can escape all strict saddle points asymptotically [Lee et al., '15].
 - Stable manifold theorem.
- Randomly initialized gradient descent can take exponential time to escape strict saddle points [Du et al., '17].

If 1) all local minima are global, and 2) are saddle points are strict, then noise-injected (stochastic) gradient descent finds a global minimum in polynomial time



What problems satisfy these two conditions

- Matrix factorization
- Matrix sensing
- Matrix completion
- Tensor factorization
- Two-layer neural network with quadratic activation

What about neural networks?

- Linear networks (neural networks with linear activations functions): **all local minima are global, but there exists saddle points that are not strict** [Kawaguchi '16].
- Non-linear neural networks with:
 - Virtually any non-linearity,
 - Even with Gaussian inputs,
 - Labels are generated by a neural network of the same architecture,**There are many bad local minima** [Safran-Shamir '18, Yun-Sra-Jadbaie '19].