

Simon OH Today 10:30A

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:

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

Handwritten notes: $f(x_t) \downarrow$, $\|\nabla f(x_t)\|_2^2$, $\downarrow$ LR

Converging to stationary points

$$X: \nabla f(X) = 0$$

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

ϵ - approximate stationary points

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

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

Sum over $t = 0, \dots, T-1$

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

$$\Rightarrow f(X_T) \leq f(X_0) - \frac{\eta}{2} \sum_{t=0}^{T-1} \|\nabla f(X_t)\|_2^2$$

$$\Rightarrow \frac{\eta}{2} \sum_{t=0}^{T-1} \|\nabla f(X_t)\|_2^2 \leq f(X_0) - f(X_T) \leq f(X_0) - \min_X f(X)$$

$$\Rightarrow T \cdot \frac{\eta}{2} \min_{0 \leq t \leq T-1} \|\nabla f(X_t)\|_2^2 \leq f(X_0) - \min_X f(X) \leq \epsilon$$

$$\Rightarrow \min_{0 \leq t \leq T-1} \|\nabla f(X_t)\|_2 \leq \sqrt{\frac{2\beta(f(X_0) - \min_X f(X))}{T}} \Rightarrow T = O\left(\frac{\beta}{\epsilon^2}\right)$$

Gradient Descent for Quadratic Functions

Opt: $x=0$

symmetric, full-rank, $\lambda_{\min}(A) > 0$

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

$$\|x_t\|_2 \leq (1 - \eta \lambda_{\min})^t \|x_0\|_2 \quad \text{to make } \|x_t\| \leq \epsilon$$

$$\|x_{t+1}\|_2 = \|x_t - \eta A x_t\|_2$$

$$= \|(I - \eta A) x_t\|_2$$

$$\leq \|I - \eta A\|_2 \cdot \|x_t\|_2$$

$$= (1 - \eta \lambda_{\min}) \cdot \|x_t\|_2$$

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

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

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

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

• Can be generalized strongly convex

$$\text{e.g. } A = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Momentum: Heavy-Ball Method (Polyak '64)

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

Problem: $\min_x f(x)$

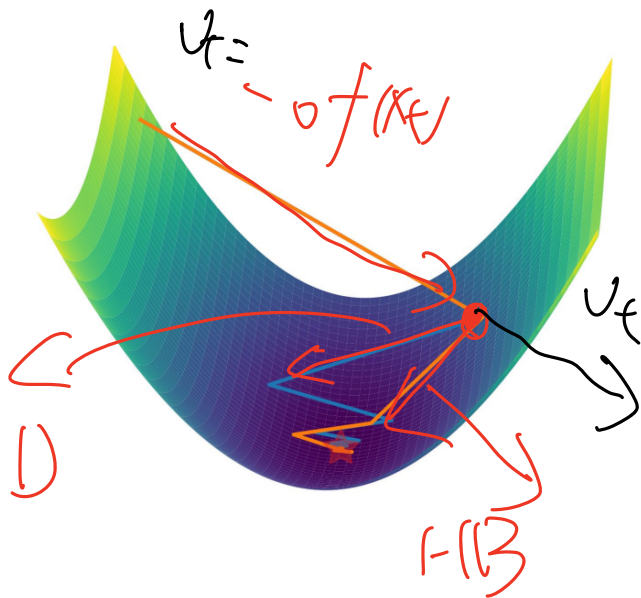
Method: $v_{t+1} = -\nabla f(x_t) + \beta v_t$
 $x_{t+1} = x_t + \eta v_{t+1}$

For quadratic function:
provable improvement

$$O(\sqrt{K} \cdot \log(\frac{1}{\epsilon}))$$

$$\text{vs. } O(K \cdot \log(\frac{1}{\epsilon}))$$

• does not work
for strongly convex



Momentum: Nesterov Acceleration (Nesterov '89)

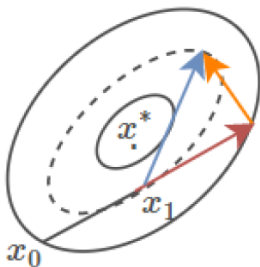
Problem: $\min_x f(x)$ *- of $f(x_t)$ for $1 \leq t \leq B$*

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

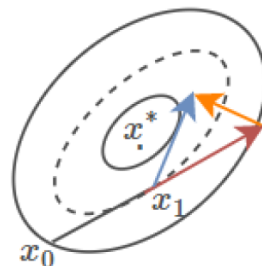
For general strongly convex functions: $O(\sqrt{L \log \frac{1}{\epsilon}})$

• use ODE

Polyak's Momentum



Nesterov Momentum



Newton's Method

2nd order method
old³

Newton's Method: $x_{t+1} = x_t - \underbrace{\eta}_{\text{step size}} \underbrace{(\nabla^2 f(x_t))^{-1} \nabla f(x_t)}_{\text{isotropic gradient}}$

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

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

$\Delta \propto -\nabla f(x)$

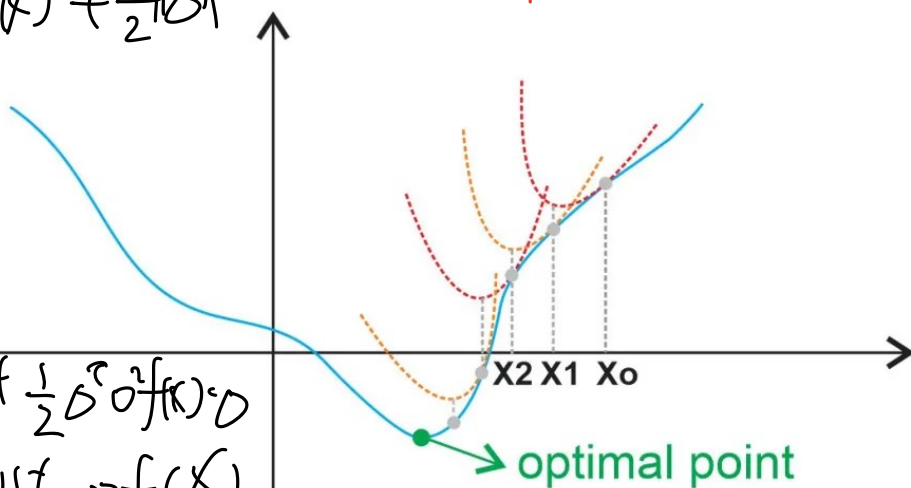
• Newton:

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

$\Rightarrow \Delta \propto -(\nabla^2 f(x))^{-1} \cdot \nabla f(x)$

$T = O\left(\log \log \left(\frac{1}{\epsilon}\right)\right)$

$\rightarrow$ isotropic gradient



AdaGrad (Duchi et al. '11)

① diagonal

② build inverse Hessian in an online

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

renormalizing

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

G_t is always diagonal

• effective learning rate: $\eta \cdot (G_t + \epsilon I)^{-1}$ small

• default hyperparameters work well

RMSProp (Hinton et al. '12)

root mean squared propagator

AdaGrad: separate learning rate for every parameter

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

RMSProp: ~~exponential weighting of gradient norms~~

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

$$0 < \beta < 1$$

$1 \leq t$
cur: $\beta^{t-1} \rightarrow 0$

AdaDelta (Zeiler '12)

RMSProp:

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

AdaDelta:

$$x_{t+1} = x_t - \eta \Delta x_t, \quad \text{unit of } x$$
$$\Delta x_t = \sqrt{u_t + \epsilon} \cdot (G_{t+1} + \epsilon I)^{-1/2} \nabla f(x_t) \quad \frac{\partial f}{\partial x} / \left(\frac{\partial f}{\partial x} \right) \propto \text{unitless}$$
$$(G_{t+1})_{ii} = \rho(G_t)_{ii} + (1 - \rho)(\nabla f(x_t)_i)^2, \quad \propto \text{unit}^2 \text{ of } x$$
$$u_{t+1} = \rho u_t + (1 - \rho) \|\Delta x_t\|_2^2$$

$$GD: \quad \Delta x \propto \frac{\partial f}{\partial x} \propto \frac{1}{\text{unit of } x}$$

$$\text{Newton:} \quad \Delta x \propto \frac{\partial f}{\partial x} / \left(\frac{\partial^2 f}{\partial x^2} \right) \propto \text{unit of } x$$

Adam (Kingma & Ba '14)

Momentum:

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

RMSProp: exponential weighting of gradient norms

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

Adam:

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

Default choice nowadays.

$\Rightarrow$ convex problem s.t. Adam does not converge

Other Optimizers

- AdamW
- NAdam
- RAdam
- GGT
- K-FAC
- ...

adaptif

Are these actually useful

Heavy-bull

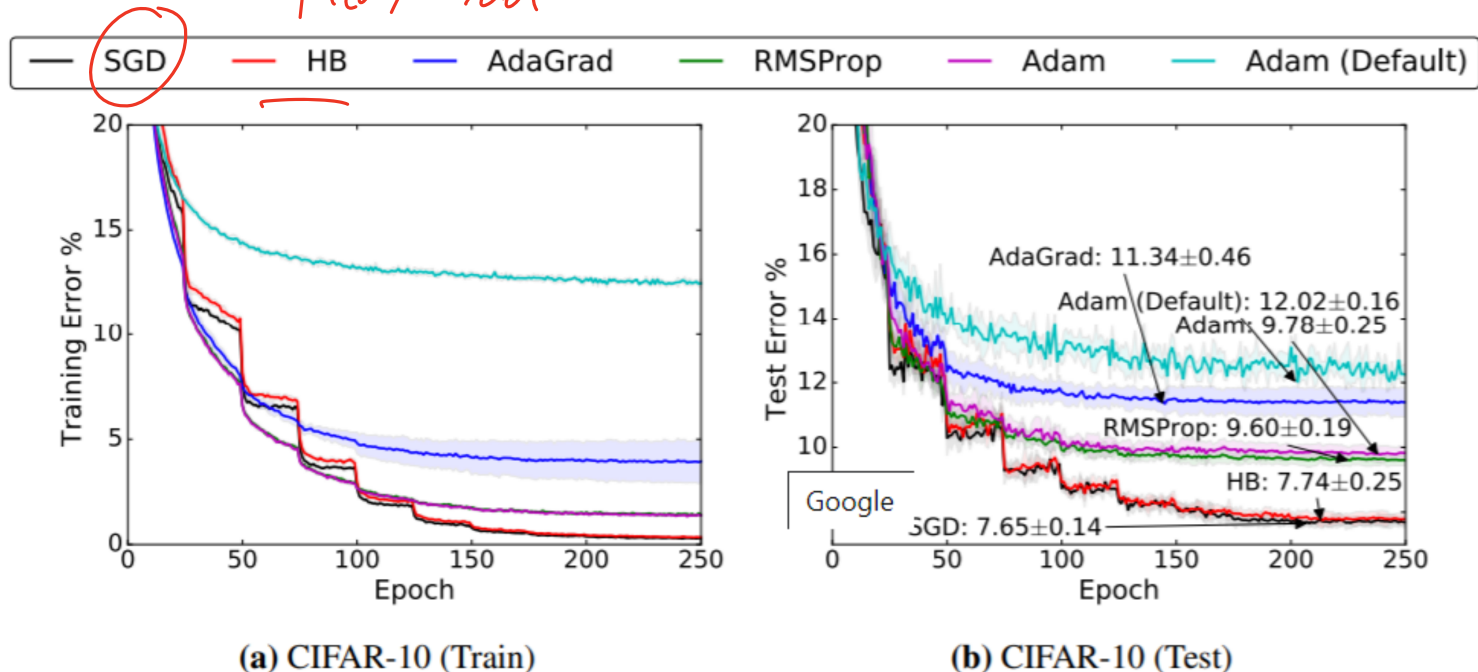


Figure 1: Training (left) and top-1 test error (right) on CIFAR-10. The annotations indicate where the best performance is attained for each method. The shading represents $\pm$ one standard deviation computed across five runs from random initial starting points. In all cases, adaptive methods are performing worse on both train and test than non-adaptive methods.

Wilson, Roelofs, Stern, Srebro, Recht '18

Important Techniques in Neural Network Training



Gradient Explosion / Vanishing

$\sigma: \text{ReLU}$

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

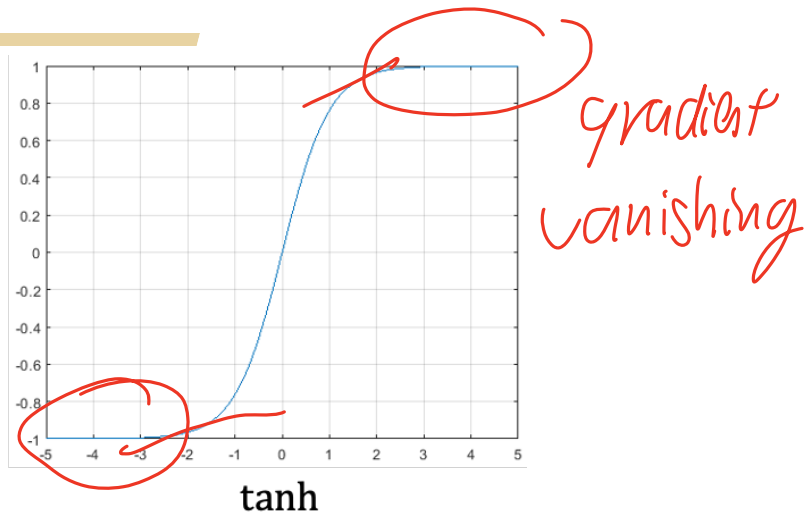
$$f(x, w_1, \dots, w_{H+1}) = w_{H+1} \sigma(w_H \dots \sigma(w_1 x) \dots)$$

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

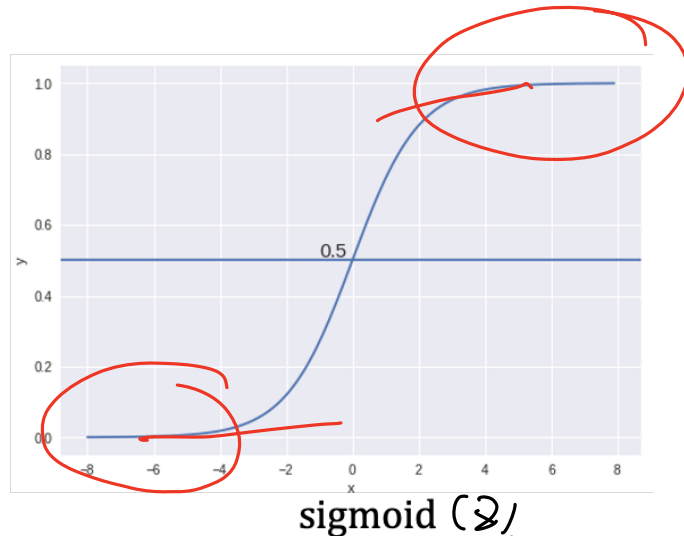
• magnitude $\uparrow$ gradient exp large

• magnitude $\downarrow$ / subspace does align
gradient exp small

Activation Functions

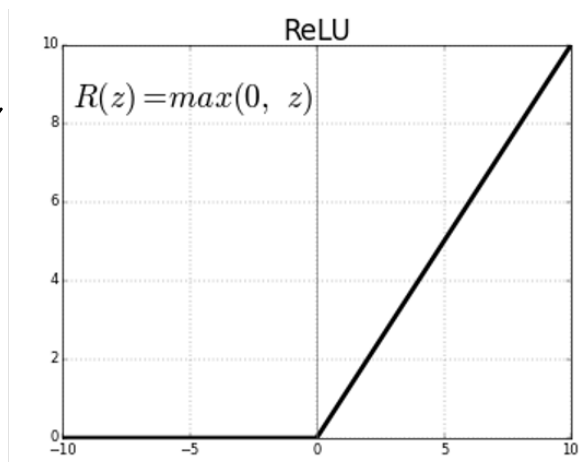


$$\tanh(z) = \frac{e^{\sigma(z)} - e^{-\sigma(z)}}{e^{\sigma(z)} + e^{-\sigma(z)}}$$



$$\text{sigmoid}(z) = \frac{1}{1 + e^{-z}}$$

$$\sigma'(z) = 0 \text{ or } 1$$

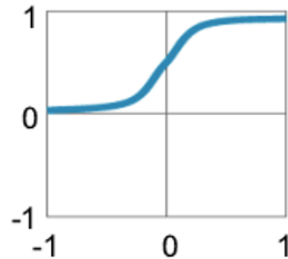


Rectified Linear Unit

Activation Function

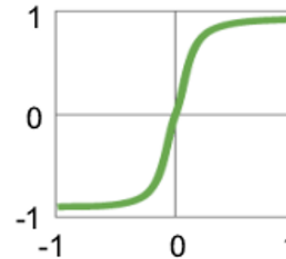
Traditional
Non-Linear
Activation
Functions

Sigmoid



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

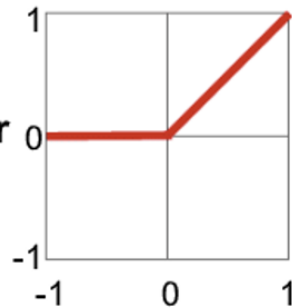
Hyperbolic Tangent



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

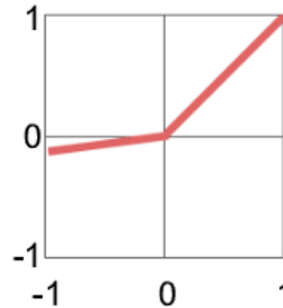
Modern
Non-Linear
Activation
Functions

Rectified Linear Unit
(ReLU)



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

Leaky ReLU

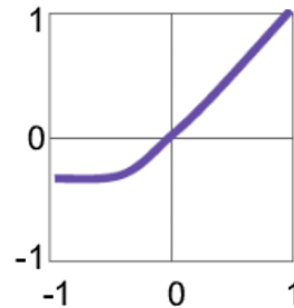


$$y = \max(\alpha x, x)$$

α = small const. (e.g. 0.1)

$\alpha \approx 0.1$

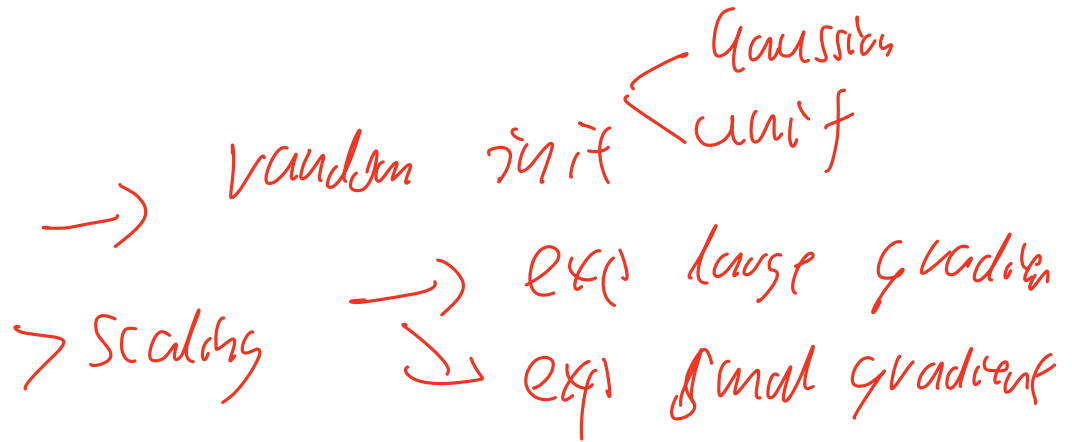
Exponential LU



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

Initialization

- Zero-initialization
- Large initialization
- Small initialization



- Design principles:
 - Zero activation mean
 - Activation variance remains same across layers

Xavier Initialization (Glorot & Bengio, '10)

- $W_{ij}^{(h)} \sim \text{Unif} \left[-\frac{\sqrt{6}}{\sqrt{d_h + d_{h+1}}}, \frac{\sqrt{6}}{\sqrt{d_h + d_{h+1}}} \right]$

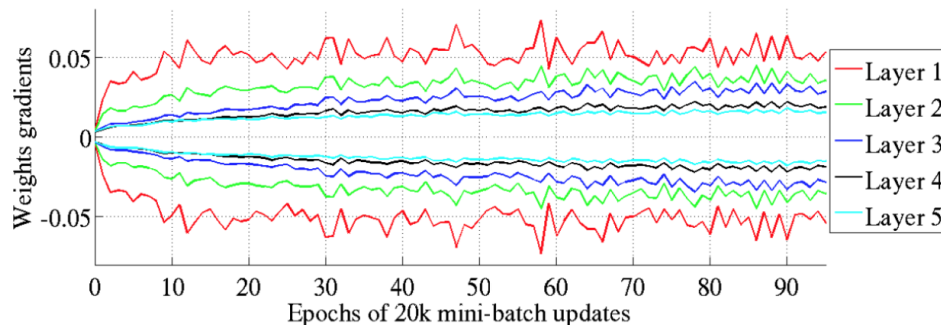
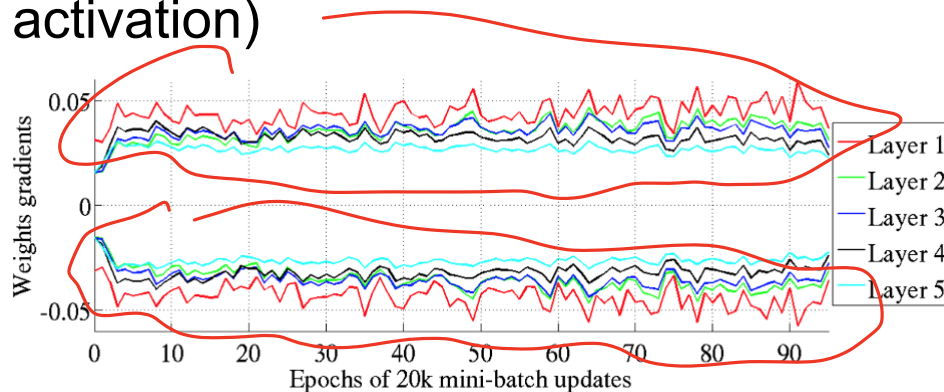
- $b^{(h)} = 0$

- Experiments (tanh activation)

$$\begin{aligned} \text{Var}(W_{ij}^{(h)}) &= \frac{1}{12} \cdot \left(\frac{2\sqrt{6}}{\sqrt{d_h + d_{h+1}}} \right)^2 \\ &= \frac{2}{d_h + d_{h+1}} \end{aligned}$$

$$\text{unif} \left(-\frac{1}{\sqrt{d_h}}, \frac{1}{\sqrt{d_h}} \right)$$

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



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

