

Lecture 22

Training neural networks

Minimize $J(\vec{\theta}) = - \sum_{n=1}^N \log p(\vec{y}_n | \vec{x}_n; \vec{\theta})$

$\uparrow$
 NN weights + biases

+ regularizer [optionally]
 $\rightarrow$ no convexity guarantees $\leftarrow$

1. Vanishing and exploding gradients

Consider $\vec{g}_l = \frac{\partial J}{\partial \vec{z}_l} = \frac{\partial J}{\partial \vec{z}_{l+1}} \underbrace{\frac{\partial \vec{z}_{l+1}}{\partial \vec{z}_l}}_{\text{activations in layer } l} = J_l \vec{g}_{l+1}$

backward pass " J_l , Jacobian matrix

Thus $\vec{g}_l = \underbrace{J_l J_{l+1} \dots J_{L-1}}_{\text{may amplify or decrease the grad depending on the eigenvalues of the Jacobian matrices}} \underbrace{\vec{g}_L}_{\text{grad at the final layer}}$

$|\lambda_{\max}^{(i)}| = \text{spectral radius of } J_i$

$\rightarrow$ Exploding grad $\Rightarrow$ grad clipping:

$$\vec{g} \rightarrow \vec{g}' = \min(1, \frac{c}{|\vec{g}|}) \vec{g},$$

$c = \text{hyperprm.}$

$\rightarrow$ Vanishing grad $\Rightarrow$ modify activation functions or NN architecture (B)

(A)

(A) Use non-saturating activation functions such as ReLU:

$$\text{ReLU}(a) = a \mathbb{I}(a > 0)$$

However, it's possible to have "the dead ReLU problem":

some components of $\vec{a}_e = W \vec{z}_{e-1}$ may become negative $\Rightarrow$ the corresp. components of $\vec{z}_e = \mathcal{S}(\vec{a}_e)$ will be set to 0, and the units will stay off.

This problem can be solved by using the leaky ReLU:

$$\text{LReLU}(a; \alpha) = \max(\alpha a, a),$$

neg. inputs $\downarrow$ $\swarrow$ pos. inputs
 $0 < \alpha < 1$ is a hyperparam that can be learned

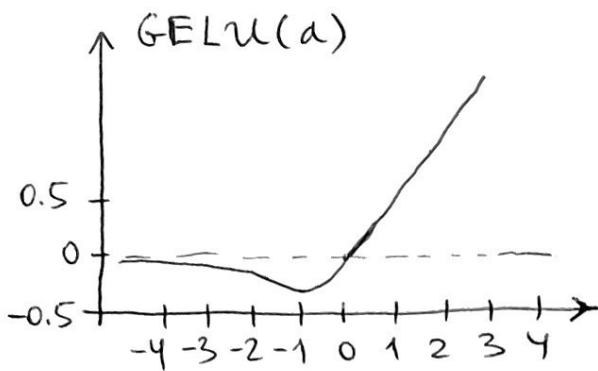
Another choice is ELU:

$$\text{ELU}(a; \alpha) = \begin{cases} \alpha(e^a - 1) & a \leq 0 \\ a & a > 0 \end{cases}$$

$\nwarrow$ smooth f'n

Finally, GELU (gaussian error linear unit) is popular:

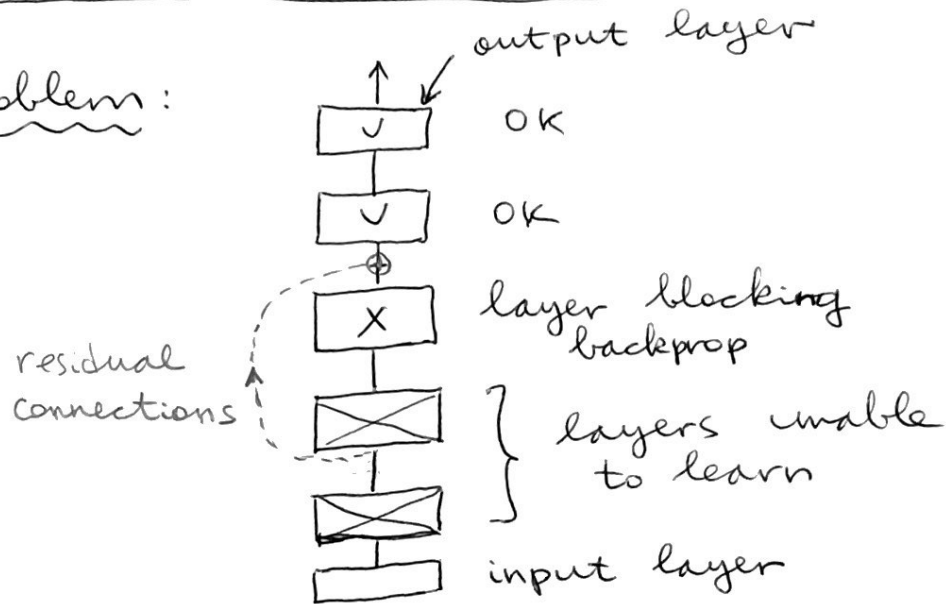
$$\text{GELU}(a) = a \Phi(a), \text{ where}$$
$$\Phi(a) = \Pr(\mathcal{N}(0, 1) \leq a) = \frac{1}{2} \left(1 + \text{erf}\left(\frac{a}{\sqrt{2}}\right) \right)$$



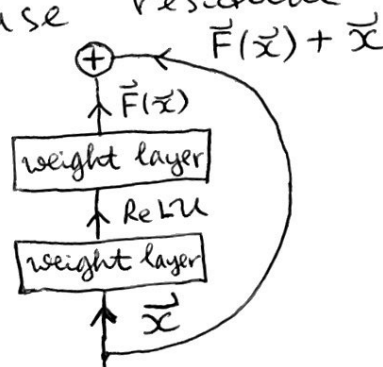
non-monotonic

(B) Residual connections (ResNet)

Problem:



Idea: use residual blocks



Later layers directly coupled to earlier layers through these "skip connections"

② Parameter initialization

Consider $a_i = \sum_{j=1}^{n_{in}} w_{ij} x_j$, where

n_{in} = # of incoming connections.

If $w_{ij} \sim \mathcal{N}(0, \sigma^2)$ and

$$\begin{cases} E[x_j] = 0, \\ \text{Var}[x_j] = \gamma^2, \end{cases} \text{ we obtain:}$$

$$E[a_i] = \sum_{j=1}^{n_{in}} E[w_{ij}] E[x_j] = 0,$$

$\uparrow$
indep. assumption

$$\text{Var}[a_i] = E[a_i^2] - \underbrace{(E[a_i])^2}_{=0} =$$

$$= E\left[\sum_{j,k=1}^{n_{in}} w_{ij} w_{ik} x_j x_k\right] = \sum_{j=1}^{n_{in}} E[w_{ij}^2 x_j^2] \quad \text{①}$$

$\uparrow$ $j \neq k$ terms do not contribute

$$\stackrel{\text{indep.}}{\downarrow} \text{①} \sum_j E[w_{ij}^2] E[x_j^2] = n_{in} \sigma^2 \gamma^2.$$

Need to set $n_{in} \sigma^2 = 1$, otherwise

$\text{Var}[a_i] \uparrow$ as $n_{in} \uparrow$.

Likewise, in the backwards pass the Var of the gradients may blow up unless $n_{out} \sigma^2 = 1$, where n_{out} = # of outgoing connections.

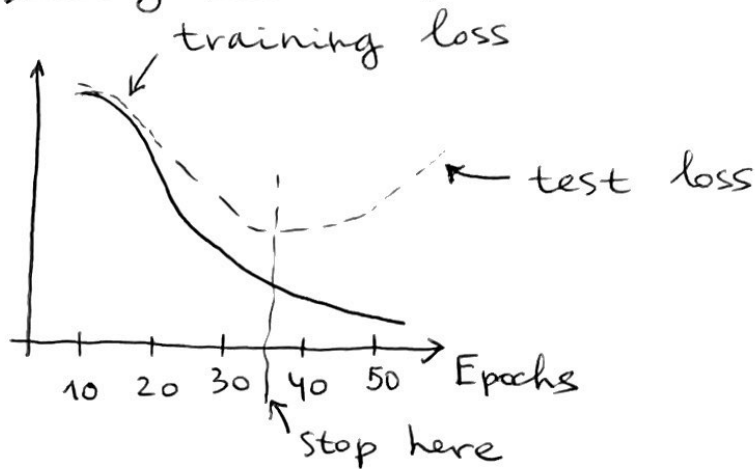
One can set

$$\sigma^2 = \frac{2}{n_{in} + n_{out}}$$

Xavier initialization

③ Regularization

① Early stopping-



② Weight decay-

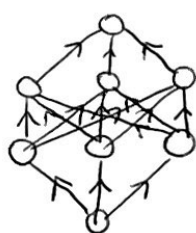
add Gaussian priors for weights and

biases :
$$\begin{cases} \mathcal{N}(\vec{w} | \vec{0}, \alpha^2 \mathbb{I}) \\ \mathcal{N}(\vec{b} | \vec{0}, \beta^2 \mathbb{I}) \end{cases}$$

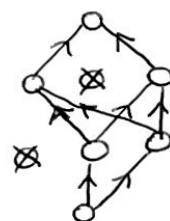
L_2 regularization;
encourages smaller
param values

③ Dropout

Turn off neurons with prob. p



$\Rightarrow$



"each neuron learns to
perform by itself, even
if other neurons are missing"

Using dropout can be viewed as estimating the noisy versions of the weights:

$$\rightarrow \theta_{ij} = w_{ij} \epsilon, \text{ where}$$

estimated during training $\epsilon = \text{Ber}(1-p)$

$$\begin{cases} i = \text{unit } i, \text{ layer } l-1 \\ j = \text{all units connected to node } i \\ \text{in layer } l \end{cases}$$

$$E[\theta_{ij}] = w_{ij}(1-p) + 0 \times p.$$

At testing, remove the dropout and rescale the weights:

$$w_{ij} = \frac{\theta_{ij}}{1-p}.$$

Finally, one can use dropout to create an ensemble of networks $\Rightarrow$ [Monte Carlo]
[dropout]

$$p(\vec{y}|\vec{x}, \mathcal{D}) = \frac{1}{S} \sum_{k=1}^S p(\vec{y}|\vec{x}; N_k)$$

$\uparrow$
 # dropout networks

$\uparrow$
 labels network k
 topology & weights