

Previous Chapter (Fundamental in ML):

$\{\bar{x}_i, y_i\}_{i=1}^N \Rightarrow P_\theta(y|\bar{x})$ Probabilistic model on input feature

Deep learning:

$\{\bar{x}_i, y_i\} : \bar{x} \xrightarrow{NN} \bar{h} \xrightarrow{P_\theta(y|\bar{h})} y$ probabilistic model on latent feature $\bar{h}$

Feed-forward Neural Network

(aka MLP: Multi-layer perceptron)

e.g. "input layer $\rightarrow$ hidden layer $\rightarrow$ hidden layer $\rightarrow$ Output layer"

Each unit: $z = \bar{w}^T \bar{x} + b$

$g(z)$: Activation function

(Non-linear) $\Rightarrow$ if all units linear, then whole model = a linear unit

Weight initialization

Xavier initialization

$w \in U[-\frac{1}{\sqrt{n}}, \frac{1}{\sqrt{n}}]$, usually used for tanh

He initialization

$w \in N[0, \frac{2}{n}]$ usually used for ReLU.

Activation function:

"step function" $\rightarrow$ not differentiable $\times$

Sigmoid, tanh, ReLU...

① Sigmoid: $g(z) = \sigma(z) = \frac{1}{1 + \exp(-z)} \Rightarrow g'(z) = g(z)[1 - g(z)]$

Pros: Easy to calculate gradients

Cons: Small gradient in most of its domain

$\Rightarrow$ Vanishing Gradients! $\Rightarrow$ Slow learning (Saturation)

$\Rightarrow$ Not Recommended in internal learning

② ReLU: Rectified Linear Units

$g(z) = \max(0, z) \Rightarrow$ Constant Gradient when $z > 0$.

fast learning!

"ReLU do not saturate"

Zero Gradient when $z < 0 \Rightarrow$ The neuron is dead (Dying ReLU)

$\Rightarrow$ To mitigate: initialize b to be a small positive value (i.e. 0.1)

$\Rightarrow$ The neuron will be initially activated

Variations: $g(z, \alpha) = \max\{0, z\} + \alpha \min\{0, z\}$.

③ Tanh: Hyperbolic Tangent

$g(z) = \tanh(z) = \frac{\exp(z) - \exp(-z)}{\exp(z) + \exp(-z)}$

larger gradient smoother still saturate

④ Softplus

$g(z) = \text{C}(z) = \log(1 + \exp(z))$

⑤ Exponential Linear Unit

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

Easy to compute derivative

⑥ Gated Linear Unit ($g(z) = z \cdot \text{"gate"}$)

SWLU $g(z) = \sigma(z)$

Swish activation $g(z) = \sigma(\beta z)$

Mish: $g(z) = z \tanh(\log(1 + \exp(z))) = \tanh(C(z))$

GELU: $g(z) = \Phi(z)$

Standard Gaussian Cdf $\Phi(z) \approx \sigma(1.7z)$

"Avoid Saturation"

SwiGLU = Swish + GLU (Used in Modern Transformer)

"More expressive"

SwiGLU(z) = Swish($wz + b$) $\cdot$ GLU($vz + c$)

Forward Propagation

Universal Approximation

1. One layer of sigmoid hidden units can approximate any well-behaved function

2. Multi-layer $\rightarrow$ less parameters.

FFN as Probabilistic Model

first-second last layer: $h = f(x)$ ("Feature extractor")

last layer: probabilistic model: $P(y|h)$ (Classification head)

$\Rightarrow$ Whole model: $P(y|x, \theta)$

$z = f(x)$, θ consists of weights in $f(\cdot)$ & parameters in $P(y|z)$

① $y \in \mathbb{R}^1 \Rightarrow P(y|z) = \mathcal{N}(y|z, \sigma^2)$ z scalar.

Linear Gaussian output Unit

Per-sample Loss: $P(x|y, \theta) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(y-z)^2}$ (Since $\sigma^2 = 1$)

$L(x, y, \theta) = -\log P(x|y, \theta) \propto \frac{1}{2}(y-z)^2$

$\therefore \frac{\partial L}{\partial z} = y - z$ Error

② $y \in \{0, 1\} \Rightarrow P(y=1|z) = \sigma(z)$

Sigmoid Output Unit

$P(y|x) = \text{Ber}(y|\sigma(z))$

$L(x, y, \theta) = -[y \log \sigma(z) + (1-y) \log (1 - \sigma(z))]$

Per-sample Loss

$\frac{\partial L}{\partial z} = \sigma(z) - y$, $\sigma(z) - y = 0 \Rightarrow \begin{cases} z \gg 0, y = 1 \\ \text{or} \\ z \ll 0, y = 0 \end{cases}$

$\rightarrow$ Already has correct answer

$\Rightarrow$ Saturation will not affect the output.

③ $y \in \{1, 2, \dots, C\}$. $P(y=k) = \text{Softmax}(\dots)$

Softmax Output Unit

$P(y=k|x) = \frac{\exp(z_k)}{\sum_{j=1}^C \exp(z_j)}$

$\frac{\partial L}{\partial z_k} = P(y=k) - 1$ if $y=k$

Chain Rule: $\frac{\partial L}{\partial w} = \frac{\partial L}{\partial z} \frac{\partial z}{\partial w}$

Back Propagation

related to probabilistic model (Above)

SGD: Stochastic Gradient Descent

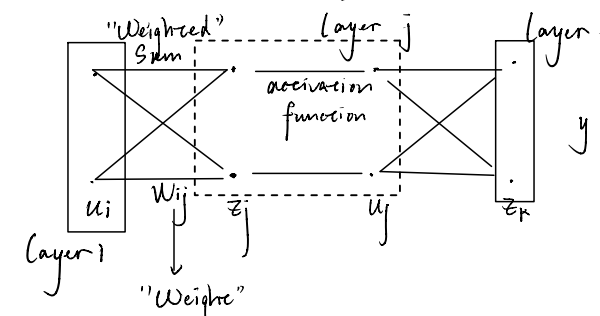
" $\theta \leftarrow \theta - \alpha \frac{1}{|\mathcal{B}|} \sum_{(x_i, y_i) \in \mathcal{B}} \nabla L(x_i, y_i, \theta)$ "

Backpropagation:

$L(\theta) = \frac{1}{N} \sum_{i=1}^N L(x_i, y_i, \theta) = -\frac{1}{N} \sum_{i=1}^N \log P(y_i|x_i, \theta)$

SGD

$\theta \leftarrow \theta - \alpha \frac{1}{|\mathcal{B}|} \sum_{(x_i, y_i) \in \mathcal{B}} \nabla L(x_i, y_i, \theta)$

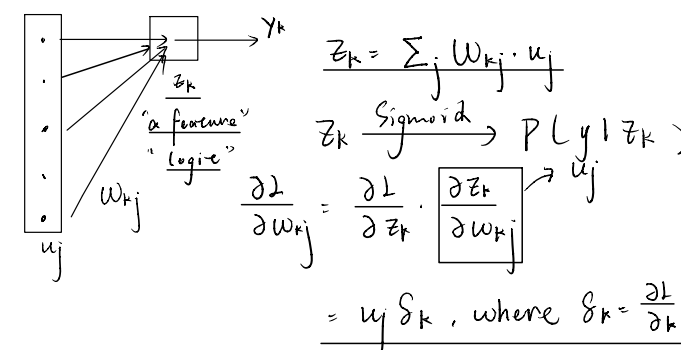


hidden layer: $u_i \xrightarrow{\sum_i w_{ki} u_i} z_k \xrightarrow{g(z_k)} u_k \dots$

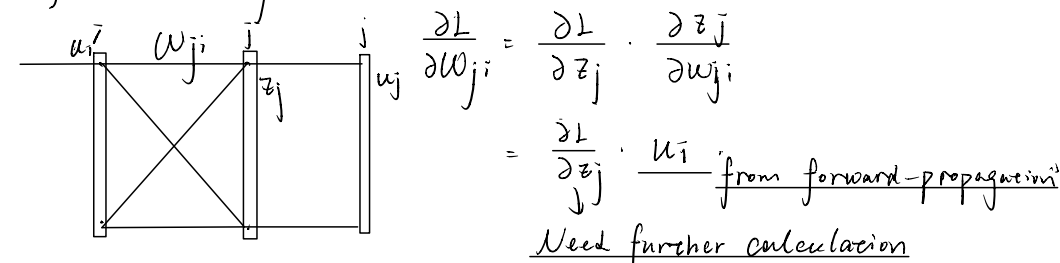
Output layer:

Chain Rule: $u_j \xrightarrow{\sum_k w_{kj} u_j} z_k \xrightarrow{\text{logit}} P(y|z_k)$

1. For output layer's units:



2. for hidden layers



Need further calculation

$\frac{\partial L}{\partial z_j} = \frac{\partial L}{\partial u_j} \frac{\partial u_j}{\partial z_j}$

since $\frac{\partial L}{\partial u_j} = \sum_k \frac{\partial L}{\partial z_k} \frac{\partial z_k}{\partial u_j} = \sum_k \frac{\partial L}{\partial z_k} w_{kj}$

$\therefore \frac{\partial L}{\partial z_j} = \sum_k \frac{\partial L}{\partial z_k} w_{kj} \cdot \frac{\partial u_j}{\partial z_j}$

$\therefore$ Back Propagation:

Output layer: $\delta_k = \frac{\partial L}{\partial z_k} \rightarrow \frac{\partial L}{\partial w_{kj}} = \frac{\partial L}{\partial z_k} \frac{\partial z_k}{\partial w_{kj}} = \delta_k \cdot \frac{\partial z_k}{\partial w_{kj}} = \delta_k \cdot u_j$

hidden layer: $\delta_j = \frac{\partial L}{\partial z_j} = \frac{\partial L}{\partial z_k} \frac{\partial z_k}{\partial u_j} \frac{\partial u_j}{\partial z_j} = \delta_k \cdot w_{kj} \cdot \frac{\partial u_j}{\partial z_j}$

activation function $g(\cdot)$
 $u_j = g(z_j)$
 $\Rightarrow \frac{\partial u_j}{\partial z_j} = g'(z_j)$

hidden layer: $\delta_j = (\sum_k \delta_k \cdot w_{kj}) \frac{\partial u_j}{\partial z_j}$

$\therefore$ Actually, we're just back-propagation " $\delta = \frac{\partial L}{\partial z}$ "