

$$\nabla_{\theta} J(\theta)$$

Preclass 08: Neural Network and

Backpropagation^a $\Leftarrow$ gradient descent, neural network
역전파.

[SCS4049] Machine Learning and Data Science

Seongsik Park (s.park@dgu.edu)

School of AI Convergence & Department of Artificial Intelligence, Dongguk University

Neural Networks and Perceptron

아랑외행

מכירת

of neuron

$\lambda \propto \frac{1}{\nu}$

LH neuron



threshold,
on/off



1/32

Computing with the brain

An engineering perspective

- Compact ←
- Energy efficient (20 watts)
- 85 billion Glial cells (power, cooling, support)
- 86 billion Neurons (soma + wires) 수백억 ←
- 69 billion Cerebellum neurons (soma + wires)
- 103 104 Connections (synapses) per neuron
- Volume = mostly wires

General computing machine?

- Slow for mathematical logic, arithmetic, etc
- Very fast for vision, speech, language, social interactions, etc
- Evolution: vision $\Rightarrow$ language $\Rightarrow$ logic 고등

컴퓨터. 고등 ←

컴퓨터 VS. 뇌
연산 >>>>
인식, 인지 <<<

Birds inspired us to fly, burdock plants inspired velcro, and countless more inventions were inspired by nature. It seems only logical, then, to look at the brain's architecture for inspiration on how to build an intelligent machine. This is the key idea that sparked artificial neural networks (ANNs). However, although planes were inspired by birds, they don't have to flap their wings. Similarly, ANNs have gradually become quite different from their biological cousins. Some researchers even argue that we should drop the biological analogy altogether (e.g., by saying "units" rather than "neurons"), lest we restrict our creativity to biologically plausible systems.

ANNs are at the very core of Deep Learning. They are versatile, powerful, and scalable, making them ideal to tackle large and highly complex Machine Learning tasks, such as classifying billions of images (e.g., Google Images), powering speech recognition services (e.g., Apple's Siri), recommending the best videos to watch to hundreds of millions of users every day (e.g., YouTube), or learning to beat the world champion at the game of Go by playing millions of games against itself (DeepMind's Alpha-Zero).

Threshold logic unit (TLU)

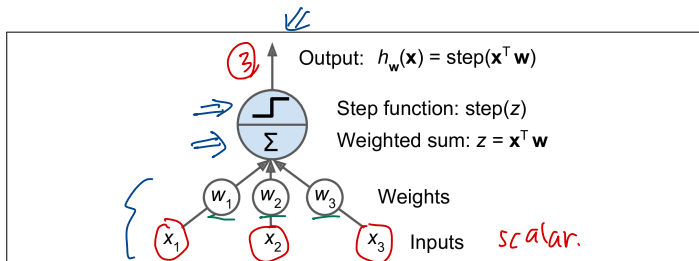


Figure 10-4. Threshold logic unit

Perceptron $\Rightarrow$ binary classifier

$$\textcircled{1} \quad \underline{z = w_1 x_1 + w_2 x_2 + \dots + x_m x_m} = \begin{bmatrix} w_1 \\ w_2 \\ \vdots \\ w_m \end{bmatrix}^T \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_m \end{bmatrix} = \mathbf{w}^T \mathbf{x} \quad (1)$$

Threshold logic unit (TLU) or linear threshold unit (LTU)

$$h_{\mathbf{w}}(\mathbf{x}) = \text{step}(z) = \textcircled{2} \underline{\text{step}(\mathbf{w}^T \mathbf{x})} \quad (2) \quad 5/32$$

Threshold logic unit

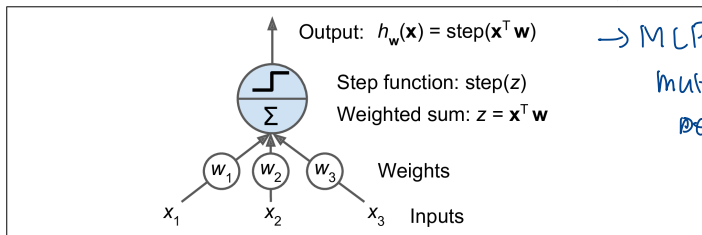


Figure 10-4. Threshold logic unit

Common step functions used in Perceptrons

$$\text{heavyside}(z) = \begin{cases} 0 & \text{if } z < 0 \\ 1 & \text{otherwise} \end{cases}$$

$$\text{sign}(z) = \begin{cases} -1 & \text{if } z < 0 \\ 0 & \text{if } z = 0 \\ +1 & \text{otherwise} \end{cases}$$

(3)

$\Rightarrow$ perceptron method.

Perceptron algorithm

Computing the outputs of a fully connected layer

$$h_{\mathbf{W}, \mathbf{b}} = \phi(\mathbf{XW} + \mathbf{b}) \quad (4)$$

Perceptron learning rule (weight update)

$$w_{i,j} \leftarrow w_{i,j} + \eta(y_j - \hat{y}_j)x_i \quad (5)$$

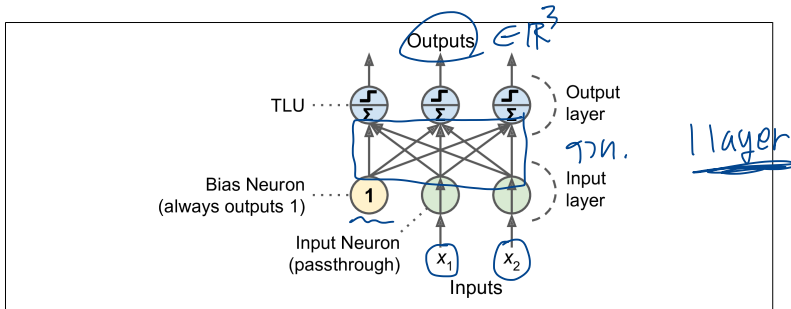


Figure 10-5. Perceptron diagram

Deep Feedforward Networks

Multi-layer perceptron algorithm

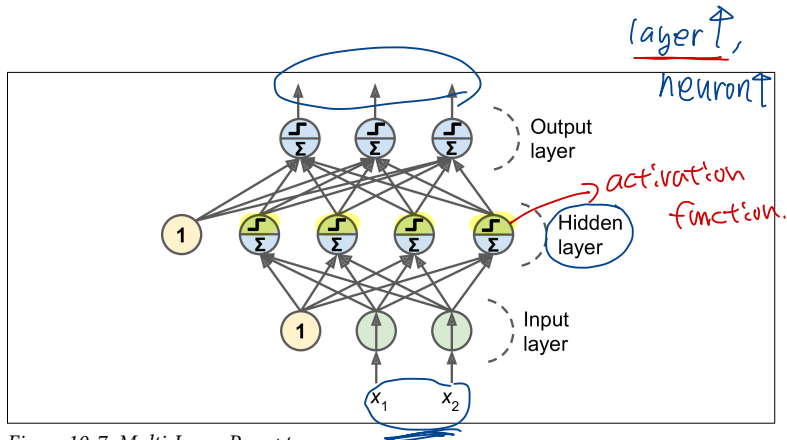


Figure 10-7. Multi-Layer Perceptron

Activation functions and their derivatives

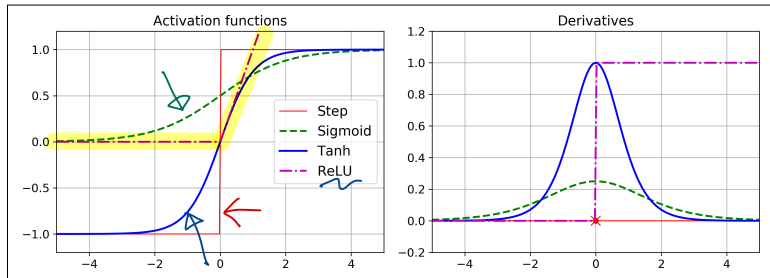


Figure 10-8. Activation functions and their derivatives

A modern MLP (including ReLU and softmax) for classification

3 classes
classification problem.

$$1 = \sum y_1 + y_2 + y_3 \rightarrow \text{probability}$$

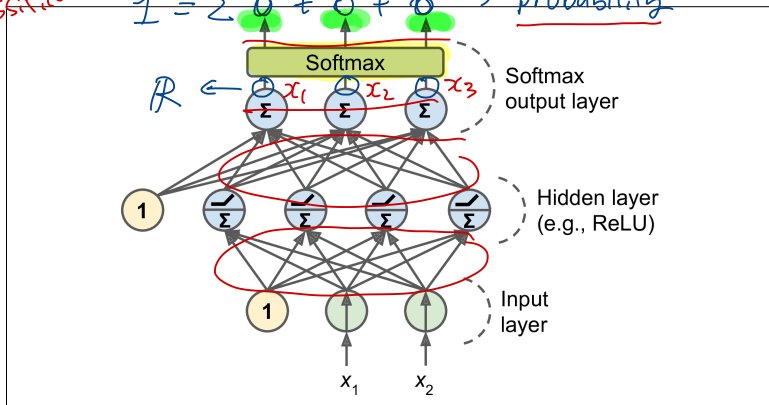


Figure 10-9. A modern MLP (including ReLU and softmax) for classification

$$y_1 = \frac{\exp(x_1)}{\exp(x_1) + \exp(x_2) + \exp(x_3)} \Leftarrow \text{softmax}$$

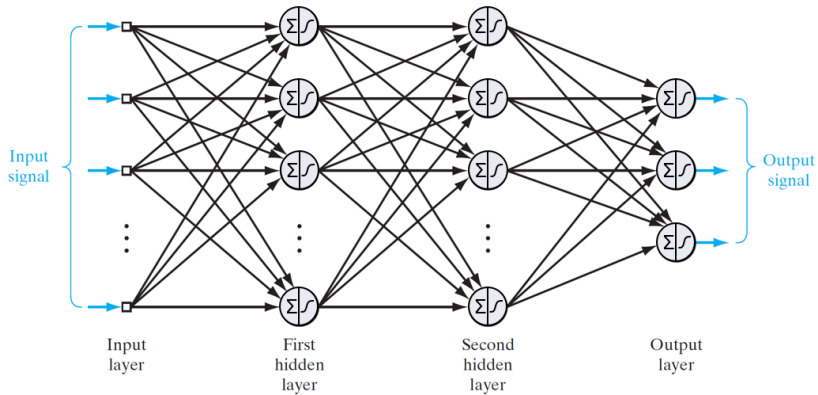
Deep feedforward networks

Deep Feedforward Networks, Multi-Layer Perceptron, or Feedforward Neural Networks

- Fully Connected Multi-layer
- Feedforward $\Rightarrow$ Feedback이 없음
(Feedback 이 있는 신경망: Recurrent Neural Networks)
- Nonlinear Activation Functions
- Rumelhart (1986) 의 Backpropagation 알고리즘 이후 집중적 연구



Deep feedforward networks



Nonlinear activation functions

Logistic sigmoid function

$$\sigma(z) = \frac{1}{1 + e^{-z}} \quad (6)$$

$$\sigma'(z) = \sigma(z)(1 - \sigma(z)) \quad (7)$$

Hyperbolic tangent function

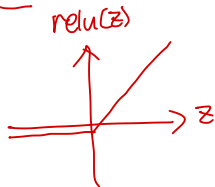
$$\varphi(z) = \tanh(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}} \quad (8)$$

$$\varphi'(z) = 2\varphi(2z) - 1 \quad (9)$$

Rectified linear unit function

$$\text{ReLU}(z) = \max(0, z) \quad (10)$$

$$\text{ReLU}'(z) = \begin{cases} 0 & \text{if } z \leq 0 \\ 1 & \text{otherwise} \end{cases} \quad (11)$$



Nonlinear activation functions

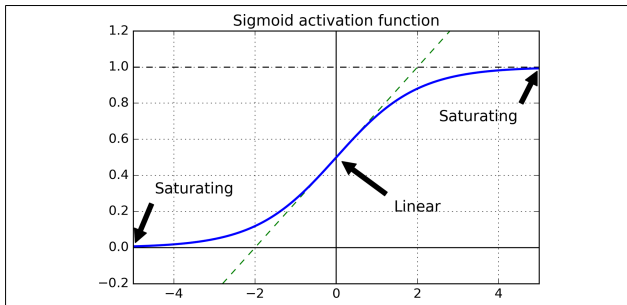


Figure 11-1. Logistic activation function saturation

Nonlinear activation functions

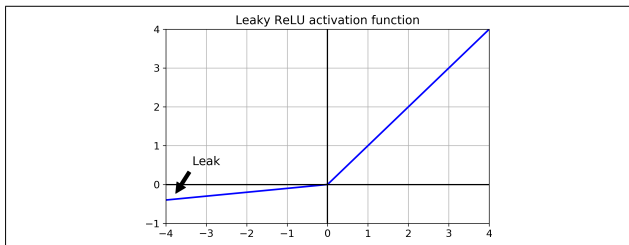
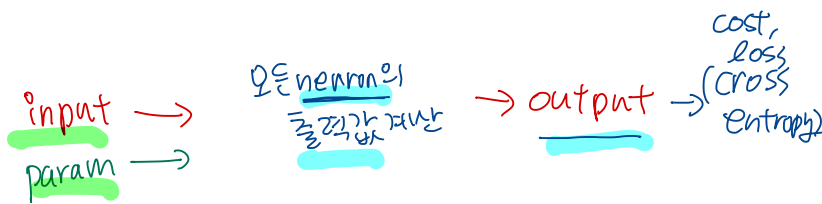
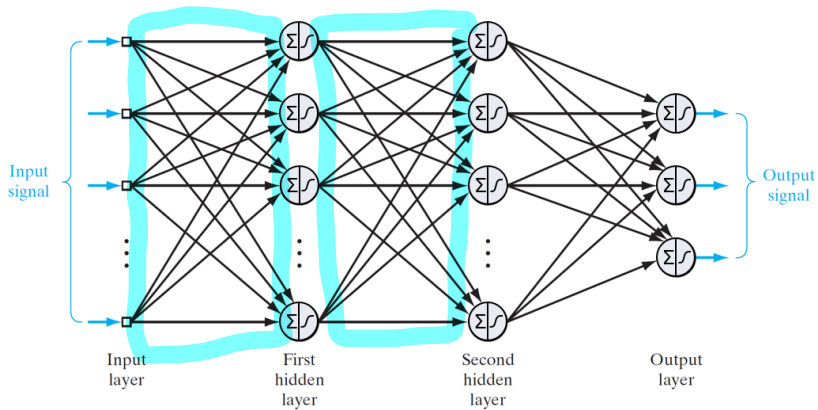


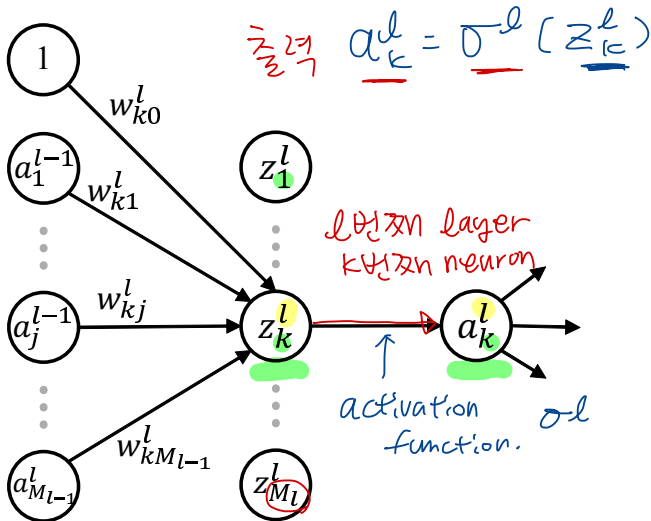
Figure 11-2. Leaky ReLU



Forward propagation

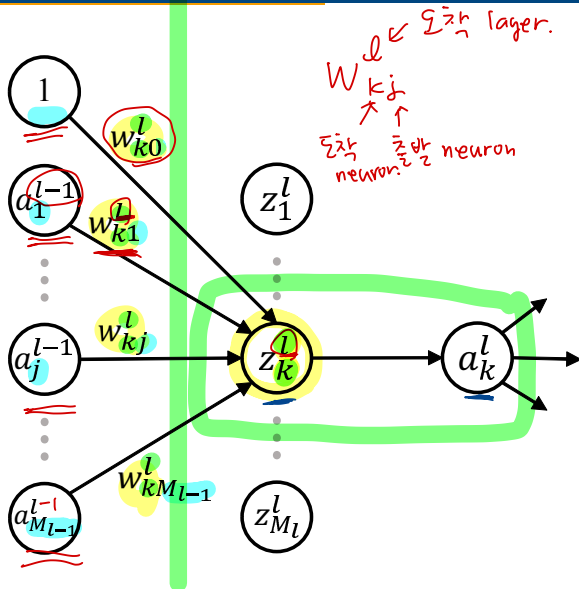
Deep feedforward networks





l번째 layer, 뉴런이 M_l 개.

Feedforward propagation



$$Z_K^l = W_{K0}^l \cdot 1 + W_{K1}^l a_1^{l-1} + W_{K2}^l a_2^{l-1} + \dots$$

$$\quad \quad \quad \underline{1 = a_0^{l-1}} \quad \quad \quad + W_{K M_{l-1}}^l a_{M_{l-1}}^{l-1}$$

$$= \sum_{\substack{j=0 \\ \text{---}}}^{M_{l-1}} W_{Kj}^l a_j^{l-1}$$

Feedforward propagation

For lth hidden layer

a:

$$\Rightarrow \underline{z}_k^l = \sum_{j=0}^{M_{l-1}} w_{kj}^l \underline{a}_j^{l-1}$$
$$\underline{\mathbf{z}}^l = \underline{\mathbf{W}}^l \underline{\mathbf{a}}^{l-1}$$

$$\underline{a}_k^l = \sigma^l(\underline{z}_k^l) \quad (12)$$

$$\underline{\mathbf{a}}^l = \sigma^l(\underline{\mathbf{z}}^l) \quad (13)$$

k neuron

- z_k^l : layer l 의 unit k 로 들어오는 activation들의 weighted sum
- w_{kj}^l : layer $(l-1)$ 의 unit j 에서 layer l 의 unit k 로 가는 가중치
- w_{k0}^l : layer l 의 unit k 에 대한 bias
- a_j^{l-1} : layer $(l-1)$ 의 unit j 에서 나오는 activation 값
- a_0^l : layer l 의 bias에 대응 ($= 1$)
- a_k^l : layer l 의 unit k 의 activation 값
- σ^l : layer l 의 activation function
- M_{l-1} : layer $(l-1)$ 의 hidden unit 수

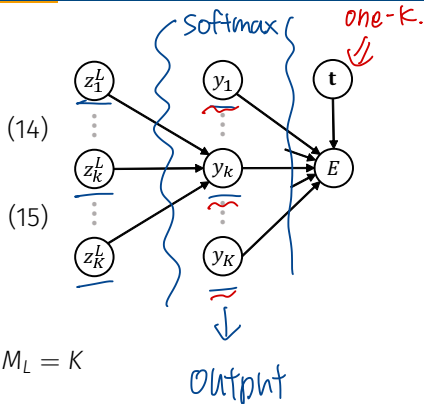
Cost function

Cost function: cross-entropy

$$E = -\frac{1}{N} \sum_{n=1}^N \sum_{j=1}^{M_L} \underline{t_j} \log \underline{y_j}$$

$$y_k = \frac{\exp(z_k)}{\sum_{j=1}^K \exp(z_j)}$$

- t_j : $\mathbf{x}$ 의 class label (target value)
- K : 출력층의 unit 수 = class의 수, $M_L = K$



$$\begin{array}{c} y \\ \hline 0.8 \\ 0.1 \\ 0.1 \\ \hline \end{array} \quad \begin{array}{c} t \\ \hline 1 \\ 0 \\ 0 \\ \hline \end{array}$$

$$- (1 \log 0.8 + 0 \log 0.1 + 0 \log 0.1)$$

Stochastic gradient and backpropagation

SGD weight update rule

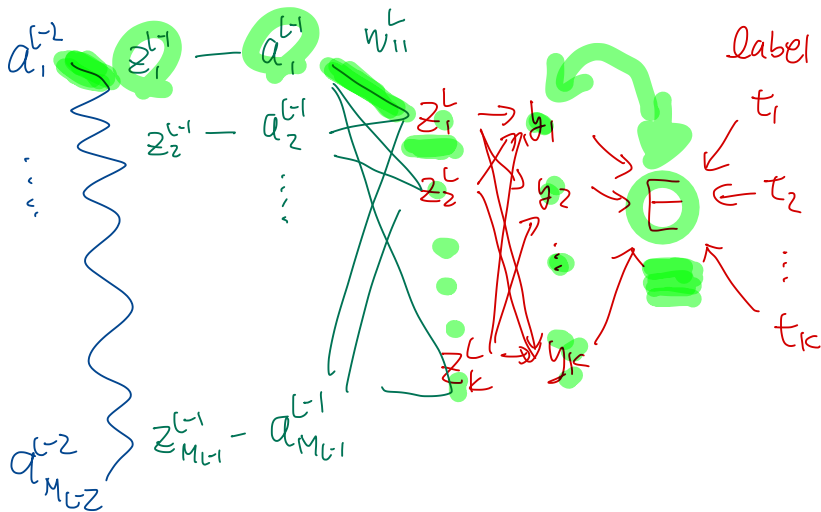
$$w_{ij}^l \leftarrow w_{ij}^l - \eta \frac{\partial E}{\partial w_{ij}^l} \quad (16)$$

여러 층을 거쳐서

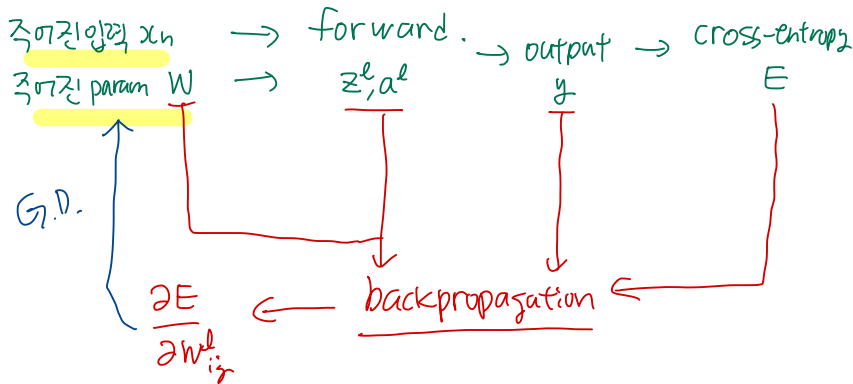
$$\mathbf{W} \leftarrow \mathbf{W} - \eta \nabla_{\mathbf{W}} E \quad (17)$$

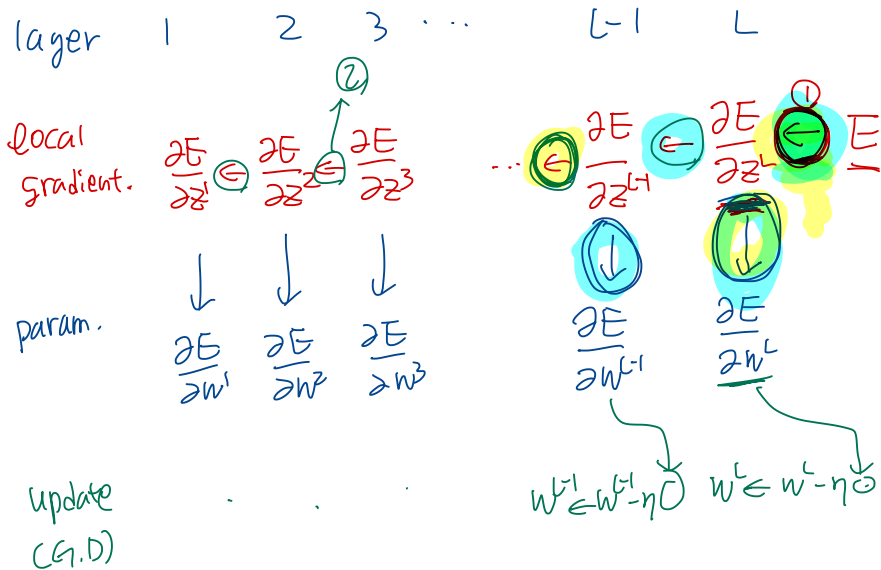
Credit assignment

- Weight 값들이 잘못 되었기 때문에 Cost 가 발생
⇒ 따라서 모든 weight에게 그 weight가 cost 발생에 기여한 만큼씩 gradient 를 배분해 주는 것
- 문제는 weight 들이 여러 층에 분산되어 있으며, 각 층은 비선형 함수이므로 배분 방법이 단순하지 않음 ⇒ backpropagation algorithm



Backpropagation





출력층부터 입력층까지 역순으로 모든 layer에서 weight w_{ij}^l 을 업데이트

$$w_{ij} \leftarrow w_{ij} - \eta \frac{\partial E}{\partial w_{ij}} \quad (18)$$

Local gradient (or error of neuron)의 정의

$$\delta_j^l \triangleq \frac{\partial E}{\partial z_j^l} \quad (19)$$

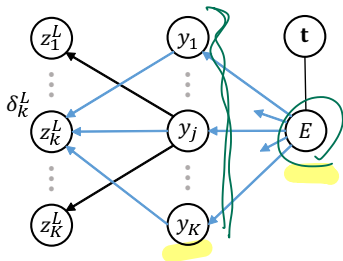
- Layer l 의 unit k 에서 error, 또는 하위층으로 분배할 gradient
- Activation a_j^l 를 기준으로 생각할 수 있으나, nonlinear 함수를 포함하므로 다루기 어려움

Backpropagation: output layer 1/7

출력층에서의 local gradient δ_k^L 계산

$$\delta_k^L \equiv \frac{\partial E}{\partial z_k^L} \equiv \sum_j \frac{\partial E}{\partial y_j} \frac{\partial y_j}{\partial z_k^L} \quad (20)$$

- Softmax 함수는 $z_1^L, z_2^L, \dots, z_K^L$ 에 의존
- 따라서 z_k^L 의 변동은 $y_1, y_2, \dots, y_K$ 에 모두 영향



Backpropagation: output layer 2/7

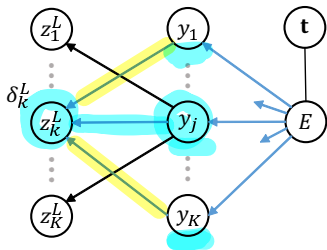
$\frac{\partial y_j}{\partial z_k^L}$ 의 계산: $k = j$ 인 경우

$$\frac{\partial y_k}{\partial z_k^L} = \frac{\partial}{\partial z_k^L} \frac{\exp(z_k)}{\sum_i \exp(z_i)}$$

$$= \frac{\exp(z_k)}{\sum_i \exp(z_i)} \left(1 - \frac{\exp(z_k)}{\sum_i \exp(z_i)} \right) \quad (21)$$

$$= y_k(1 - y_k) \quad (22)$$

$$(23)$$



$$y_1 = \frac{\exp(z_1)}{\sum \exp(z_k)}$$

$$y_1 = \frac{\exp(z_1)}{\sum \exp(z_k)}$$

$$\frac{\partial y_1}{\partial z_1} = \underbrace{\left(\frac{\partial}{\partial z_1} \exp(z_1) \right)}_{+ \exp(z_1)} \left(\frac{1}{\sum \exp(z_k)} \right)'$$

$$= \frac{\exp(z_1)}{\sum \exp(z_k)} + \frac{\exp(z_1)}{\sum \exp(z_k)} \left(- \frac{\exp(z_1)}{(\sum \exp(z_k))^2} \right)$$

$$= \frac{\exp(z_1)}{\sum \exp(z_k)} \left(1 - \frac{\exp(z_1)}{\sum \exp(z_k)} \right)$$

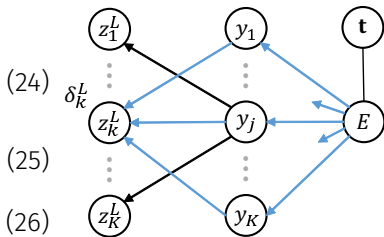
Backpropagation: output layer 3/7

$\frac{\partial y_j}{\partial z_k^L}$ 의 계산: $k \neq j$ 인 경우

$$\begin{aligned}
 \frac{\partial y_j}{\partial z_k^L} &= \frac{\partial}{\partial z_k^L} \frac{\exp(z_j)}{\sum_i \exp(z_i)} \\
 &= - \frac{\exp(z_k)}{\sum_i \exp(z_i)} \frac{\exp(z_j)}{\sum_i \exp(z_i)} \\
 &= -y_k y_j
 \end{aligned}$$

$$y_1 = \frac{\exp(z_1)}{\sum \exp(z_k)}$$

↑



From definition of local gradient

$$\delta_k^L = \frac{\partial E}{\partial z_k^L} = - \sum_j t_j \frac{\partial}{\partial z_k} \log y_j = - \sum_j t_j \frac{1}{y_j} \frac{\partial y_j}{\partial z_k} \quad (27)$$

$$= - \frac{t_k}{y_k} \frac{\partial y_k}{\partial z_k} - \sum_{j \neq k} \frac{t_j}{y_j} \frac{\partial y_j}{\partial z_k} \quad (28)$$

$$= - \frac{t_k}{y_k} y_k (1 - y_k) - \sum_{j \neq k} \frac{t_j}{y_j} (-y_k y_j) \quad (29)$$

$$= -t_k + t_k y_k + \sum_{j \neq k} t_j y_k \quad (30)$$

$$= -t_k + \sum_{j=1}^K t_j y_k \quad (31)$$

$$= -(t_k - y_k) \quad (32)$$

Backpropagation: output layer 5/7

In matrix form

$$\delta^L = \frac{\partial E}{\partial \mathbf{z}^L} = \left(\frac{\partial \mathbf{y}}{\partial \mathbf{z}^L} \right)^T \left(\frac{\partial E}{\partial \mathbf{y}} \right) \quad (33)$$

$$= \begin{bmatrix} \frac{\partial y_1}{\partial z_1} & \cdots & \frac{\partial y_j}{\partial z_1} & \cdots & \frac{\partial y_K}{\partial z_1} \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ \frac{\partial y_1}{\partial z_k} & \cdots & \frac{\partial y_j}{\partial z_k} & \cdots & \frac{\partial y_K}{\partial z_k} \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ \frac{\partial y_1}{\partial z_K} & \cdots & \frac{\partial y_j}{\partial z_K} & \cdots & \frac{\partial y_K}{\partial z_K} \end{bmatrix} \begin{bmatrix} \frac{\partial E}{\partial y_1} \\ \vdots \\ \frac{\partial E}{\partial y_j} \\ \vdots \\ \frac{\partial E}{\partial y_K} \end{bmatrix} \quad (34)$$

$$= \begin{bmatrix} \sum_j \frac{\partial E}{\partial y_j} \frac{\partial y_j}{\partial z_1} \\ \vdots \\ \sum_j \frac{\partial E}{\partial y_j} \frac{\partial y_j}{\partial z_k} \\ \vdots \\ \sum_j \frac{\partial E}{\partial y_j} \frac{\partial y_j}{\partial z_K} \end{bmatrix} = \begin{bmatrix} -(t_1 - y_1) \\ \vdots \\ -(t_k - y_k) \\ \vdots \\ -(t_K - y_K) \end{bmatrix} \quad (35)$$

With vector form

$$\mathbf{y} = \varphi(\mathbf{z}^L) \quad (36)$$

$$\delta^L = \left(\frac{\partial \mathbf{y}}{\partial \mathbf{z}^L} \right)^T \left(\frac{\partial E}{\partial \mathbf{y}} \right) \quad (37)$$

$$= (\varphi'(\mathbf{z}^L))^T \nabla_{\mathbf{y}} E \quad (38)$$

$$= -(\mathbf{t} - \mathbf{y}) \quad (39)$$


$$\frac{\partial E}{\partial \mathbf{z}^L} \leftarrow E$$

Backpropagation: output layer 7/7

$\frac{\partial E}{\partial w_{kj}^L}$ 의 계산

$$\frac{\partial E}{\partial z^L} \rightarrow \frac{\partial E}{\partial w^L}$$

$$z_k^L = \sum_{j=0}^{M_{L-1}} w_{kj}^L a_j^{L-1} \quad (40)$$

$$\frac{\partial z_k^L}{\partial w_{kj}^L} = a_j^{L-1} \quad (41)$$

$$\frac{\partial E}{\partial w_{kj}^L} = \frac{\partial E}{\partial z_k^L} \frac{\partial z_k^L}{\partial w_{kj}^L} = -(t_k - y_k) a_j^{L-1} = \delta_k^L a_j^{L-1} \quad (42)$$

그러므로 weight update rule은

$$w_{kj}^L \leftarrow w_{kj}^L - \eta \frac{\partial E}{\partial w_{kj}^L} \quad (43)$$

$$= w_{kj}^L - \eta \delta_k^L a_j^{L-1} \quad (44)$$

$$= w_{kj}^L + \eta (t_k - y_k) a_j^{L-1} \quad (45)$$

$$(46) \quad \text{28/32}$$

Backpropagation: local gradient propagation

From δ_j^{l+1} to δ_j^l

$$\delta_k^l = \frac{\partial E}{\partial z_k^l} = \sum_j \frac{\partial E}{\partial z_j^{l+1}} \frac{\partial z_j^{l+1}}{\partial z_k^l} = \sum_j \frac{\partial z_j^{l+1}}{\partial z_k^l} \delta_j^{l+1} \quad (47)$$

$$z_j^{l+1} = \sum_{i=0} w_{ji}^{l+1} a_i^l = \sum_{i=0} w_{ji}^{l+1} \sigma^l(z_i^l) \quad (48)$$

$$\frac{\partial z_j^{l+1}}{\partial z_k^l} = w_{jk}^{l+1} \sigma'^l(z_k^l) \quad (49)$$

그러므로

$$\delta_k^l = \sigma'^l(z_k^l) \sum_j w_{jk}^{l+1} \delta_j^{l+1} \quad (50)$$

Backpropagation: weight in layer l

$\frac{\partial E}{\partial w_{kj}^l}$ 의 계산

$$z_k^l = \sum_{j=0}^{M_{l-1}} w_{kj}^l a_j^{l-1} \quad (51)$$

$$\frac{\partial z_k^l}{\partial w_{kj}^l} = a_j^{l-1} \quad (52)$$

$$\frac{\partial E}{\partial w_{kj}^l} = \frac{\partial E}{\partial z_k^l} \frac{\partial z_k^l}{\partial w_{kj}^l} = \delta_k^l a_j^{l-1} \quad (53)$$

그러므로 weight update rule은

$$w_{kj}^l \leftarrow w_{kj}^l - \eta \frac{\partial E}{\partial w_{kj}^l} \quad (54)$$

$$= w_{kj}^l - \eta \delta_k^l a_j^{l-1} \quad (55)$$

Backpropagation: summary

1. Forward propagation: 각 층 $l = 2, 3, \dots, L$ 에 대해서 다음을 계산

$$\mathbf{z}^l = \mathbf{W}^l \mathbf{a}^{l-1} \quad \mathbf{a}^l = \sigma^l(\mathbf{z}^l) \quad (56)$$

2. Output error δ^L : 출력층에서 비용함수를 이용해 local gradient 계산

$$\delta^L = -(\mathbf{t} - \mathbf{y}) \quad (57)$$

3. Backpropagate the error: 출력층부터 거꾸로 다음 식을 계산

$$\delta_k^l = \sigma^{l'}(z_k^l) \sum_j w_{jk}^{l+1} \delta_j^{l+1} \quad (58)$$

4. Output the gradient: 각 weight에 대해 gradient 계산

$$\frac{\partial E}{\partial w_{kj}^l} = \delta_k^l a_j^{l-1} \quad (59)$$

5. Update weight:

$$w_{kj}^l \leftarrow w_{kj}^l - \eta \delta_k^l a_j^{l-1} \quad (60)$$

Backpropagation: summary

