

Lecture Fys-
Stk3155/4155,
October 12, 2023

NNs (Neural Networks) part 1

- Universal approximation theorem

- Components of NNs

- (i) cost/objective function

- (ii) Architecture of NN
= model

- a) number of hidden layers

- b) number of nodes in hidden layer

c) type of activation function $\sigma(z)$

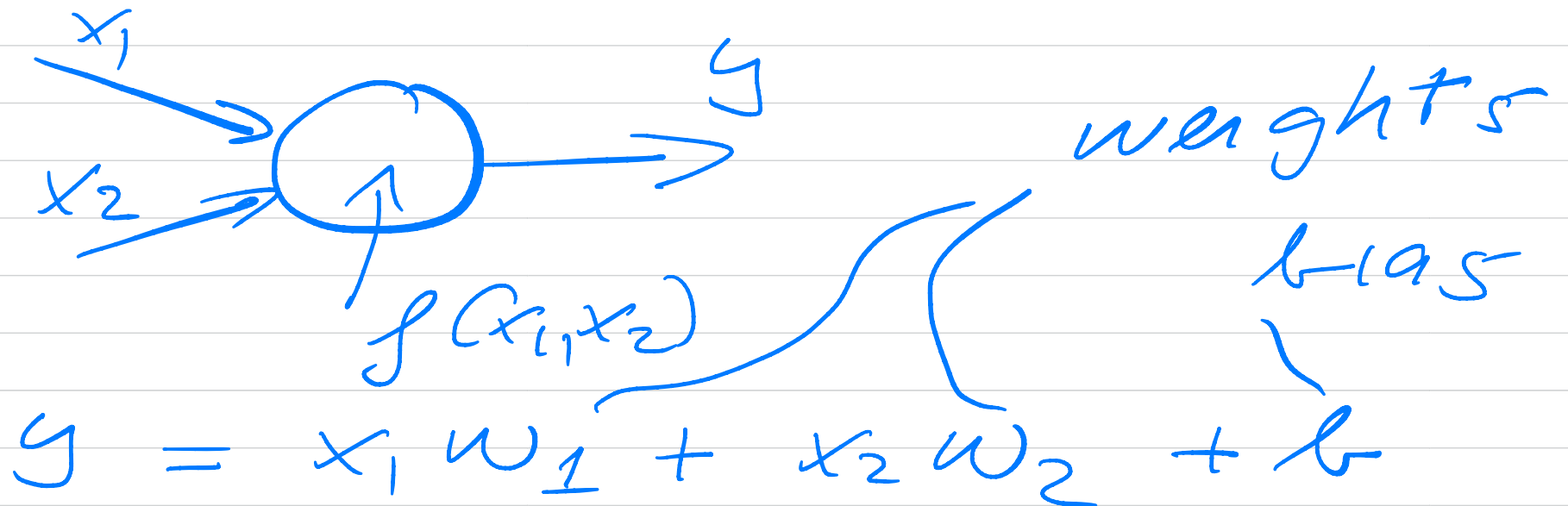
- Backpropagation for setting up gradients for training.
- Gradient methods for training (GD, SGD...)
- Hyperparameters, λ in $\lambda \sum_j \beta_j^2$ etc, learning rates etc etc

Graphical representation

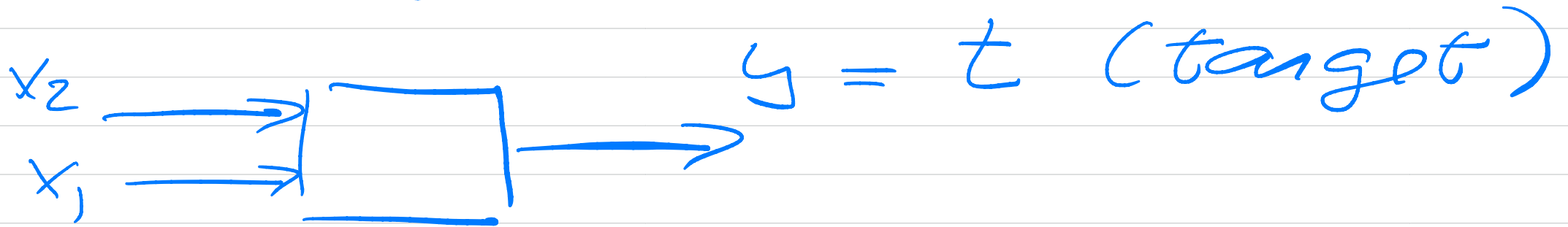
Simple network with one node



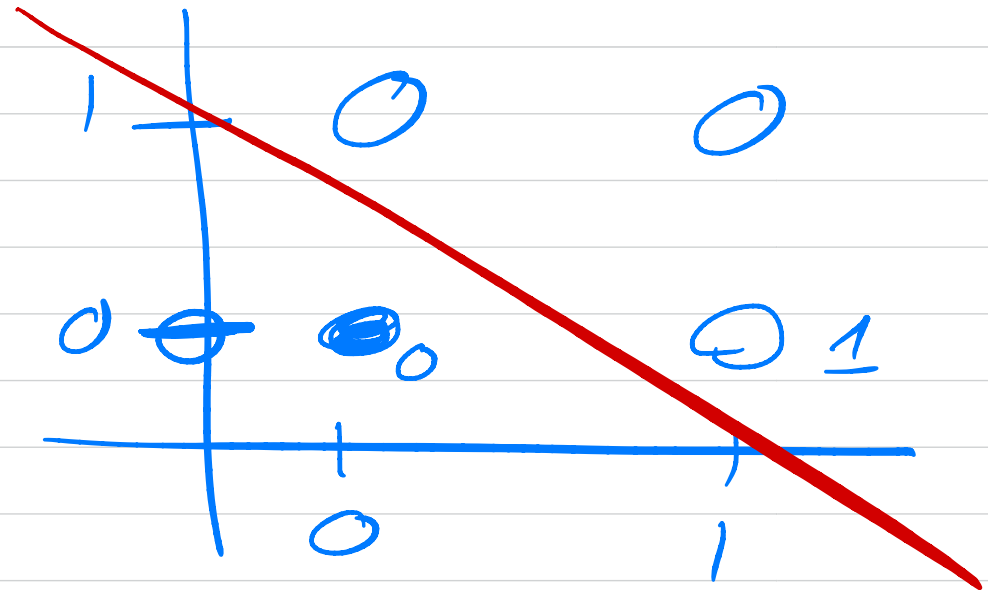
Example



OR-gate



x_1	x_2	t
0	0	0
0	1	<u>1</u>
1	0	1
1	1	<u>1</u>



$$x_1 w_1 + x_2 w_2 + b = y$$

$$w^T = [w_1, w_2]$$

$$x^T = [x_1, x_2]$$

$$y = x^T w + b$$

$$x^T = [0 \ 0], [0 \ 1], [1 \ 0], [1 \ 1]$$

$$\Theta = \{w, b\} \quad \Theta = \begin{bmatrix} b \\ w_1 \\ w_2 \end{bmatrix}$$

$$X = \begin{bmatrix} 1 & 0 & 0 \\ 1 & 0 & 1 \\ 1 & 1 & 0 \\ 1 & 1 & 1 \end{bmatrix}$$

$$\text{if } y < \frac{1}{2}$$

$$y = 0$$

$$\text{if } y > 1/2$$

$$y = 1$$

$$X^T X = \begin{bmatrix} 4 & 2 & 2 \\ 2 & 2 & 1 \\ 2 & 1 & 2 \end{bmatrix}$$

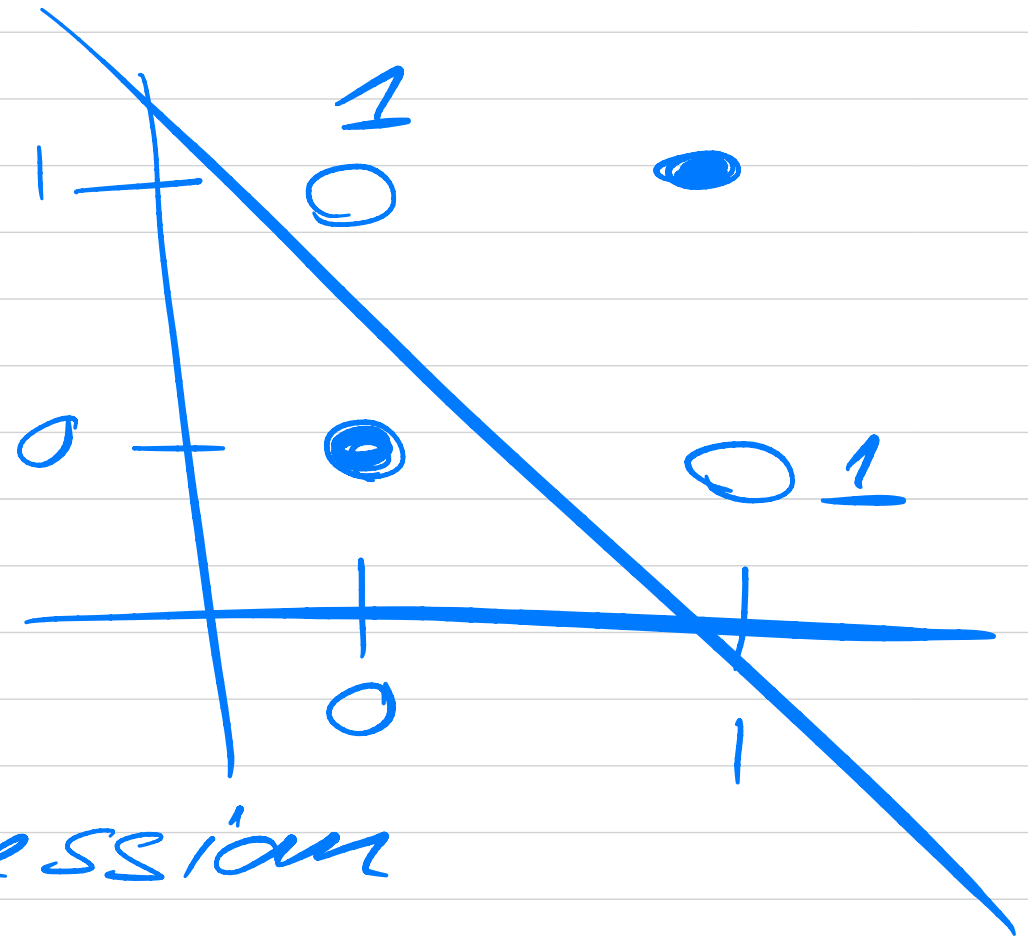
$$t = \begin{bmatrix} 0 \\ 1 \\ 1 \end{bmatrix}$$

$$\hat{\theta} = (X^T X)^{-1} X^T y$$

$$\hat{\theta} = [1/4, 1/2, 1/2] \quad \tilde{y}^T = y^T = [1/4, 3/4, 3/4, 1]$$

XOR

x_1	x_2	t
0	0	0
0	1	1
1	0	1
1	1	0



Linear regression

$$y = \left[\frac{1}{2} \frac{1}{2} \frac{1}{2} \frac{1}{2} \right]$$

$$t = [0, 1, 1, 0]$$

Universal approximation theorem
(Cybenko 1989, Hornik 1991)

Given function $F \in [0, 1]^d$
and $\varepsilon > 0$, there exists/is
a one hidden-layer NN

$f(x; \Theta)$ of the form with

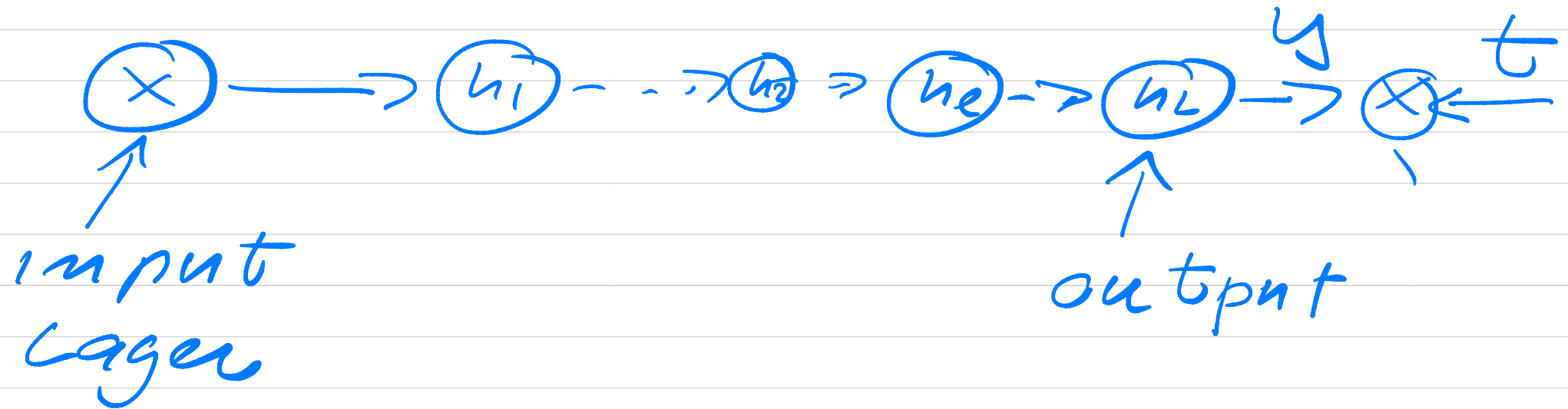
$$\Theta = \{w, b\} \quad w \in \mathbb{R}^{m \times n}$$

and $b \in \mathbb{R}^m$ for

$$\text{where } |f(x; \Theta) - F(x)| < \varepsilon$$

for all $x \in [0, 1]^d$

Back propagation Algo



- Feed Forward stage
- Back propagation stage

$$\Theta = \{w, b\}$$

Input layer



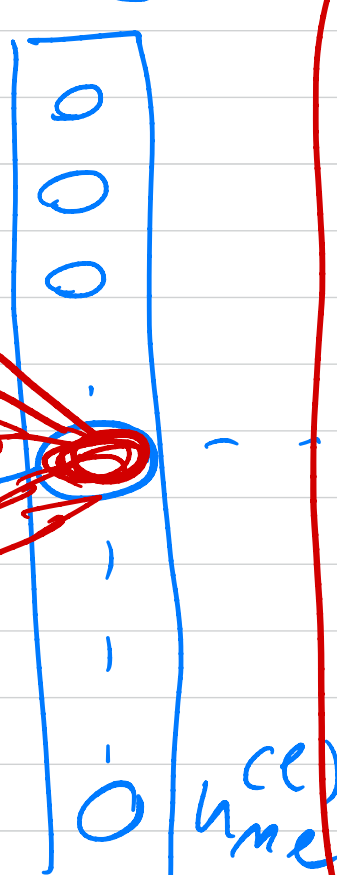
$l=1$



$l-1$

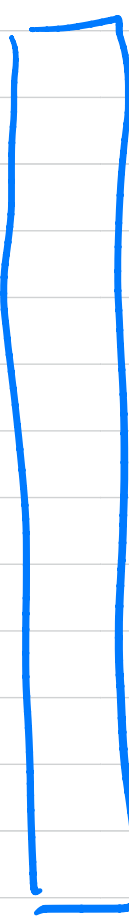


l



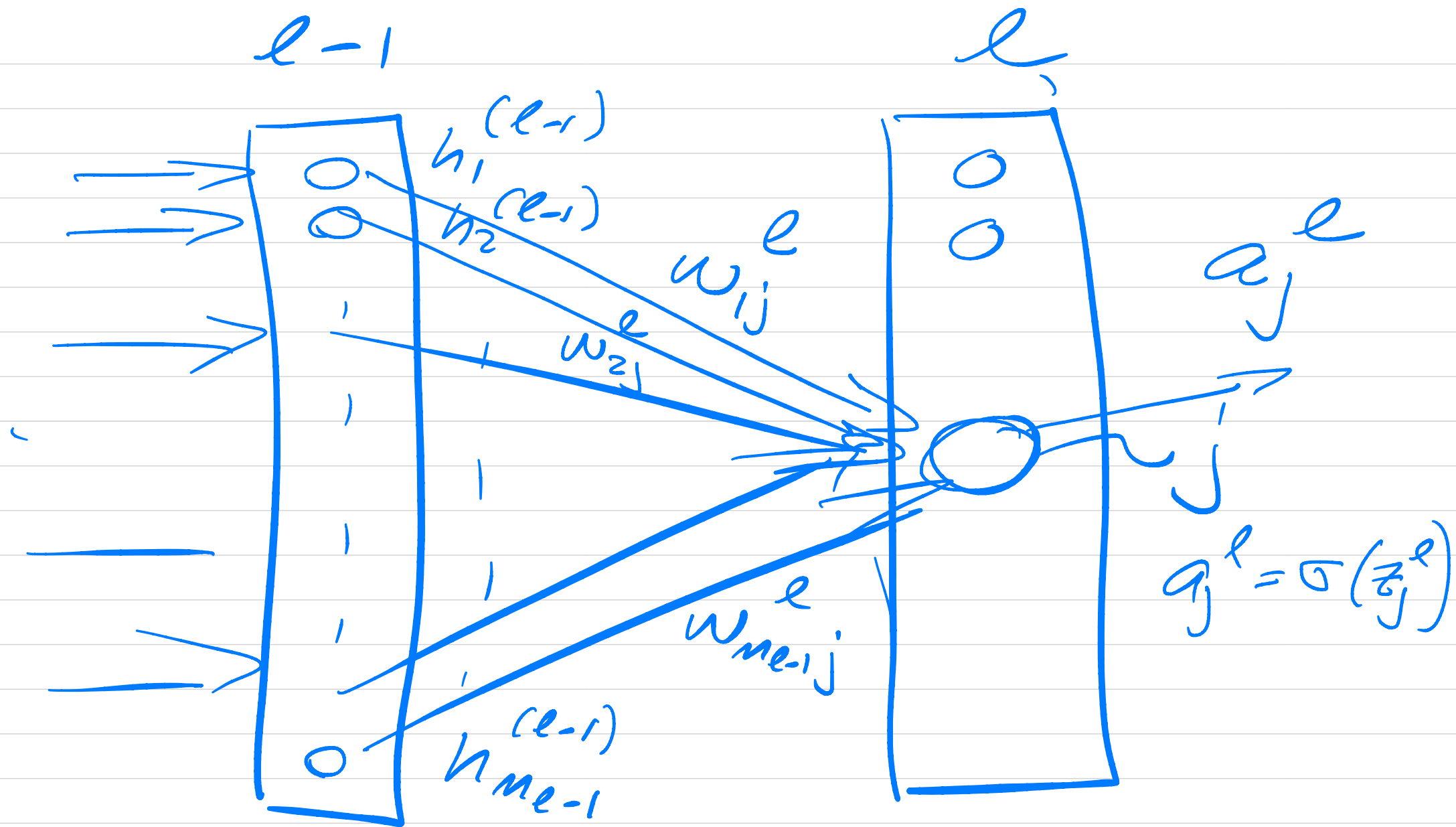
Output

L



511

target

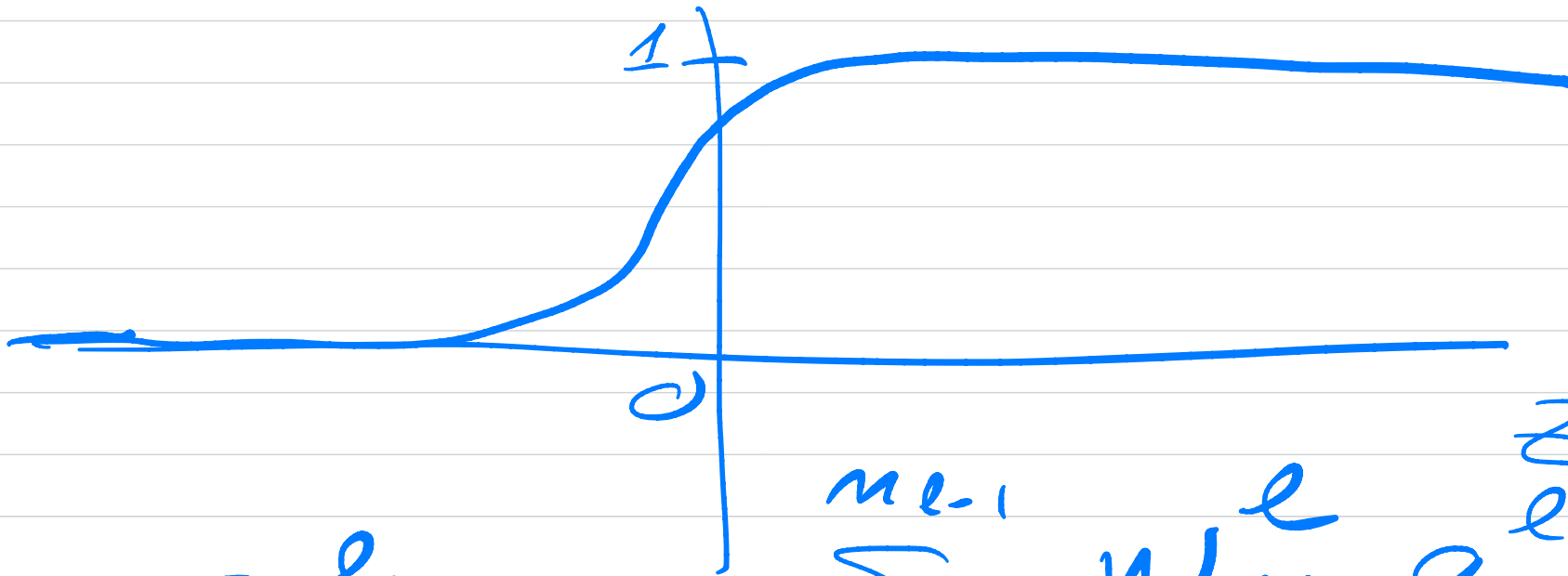


input to j in layer l
 z_j^l

$$a_j^l = \sigma(z_j^l)$$

Sigmoid is much used

$$\sigma(z_j^l) = \frac{1}{1 + e^{-z_j^l}}$$



$$z_j^l = \sum_{i=1}^{n_{l-1}} w_{ij}^l a_i^{l-1} + b_j^l$$

$$z^l = [W^{(e)}]^T a^{(l-1)} + b^{(e)}$$

W is a matrix containing the unknown parameters to train.

b is so-called bias

Need a cost function.

Here
$$C(\theta) = \frac{1}{2} \sum_{i=1}^{n_L} (y_i - t_i)^2$$

$$z_j^l = \sum_{i=1}^{m_{l-1}} w_{ij}^l a_i^{l-1} + b_j^l$$

$$a_i^{l-1} = \sigma(z_i^{l-1})$$

$$\sigma(z) = \frac{1}{1 + e^{-z}} \quad \left| \begin{array}{l} \text{in examples} \\ \text{here} \end{array} \right.$$

Some intermediate steps

$$\frac{\partial z_j^l}{\partial w_{ij}^l} = a_i^{l-1} \wedge \frac{\partial z_j^l}{\partial a_i^{l-1}} = w_{ij}^l$$

$$\frac{\partial a_j^l}{\partial z_j^l} = \frac{\sigma(z_j^l)(1 - \sigma(z_j^l))}{a_j^l(1 - a_j^l)}$$

we start at $l = L$ (output)

$$C(\theta) = \frac{1}{2} \sum_i (a_{iL} - t_i)^2$$

$$\frac{\partial C}{\partial w_{jk}^L} = (a_j^L - t_j) \frac{\partial a_j^L}{\partial w_{jk}^L}$$

$$\frac{\partial a_j^L}{\partial w_{jk}^L} = \frac{\partial a_j^L}{\partial z_j^L} \frac{\partial z_j^L}{\partial w_{jk}^L} = a_j^L (1 - a_j^L) a_k^{L-1}$$

$$\frac{\partial C}{\partial w_{jk}^L} = (a_j^L - t_j) \underbrace{a_j^L (1 - a_j^L)}_{\text{from specific activation function}} a_k^{L-1}$$

from
specific
activation
function
 $\frac{d \sigma(z)}{dz}$

$$\delta_j^L = \underbrace{a_j^L (1 - a_j^L)}_{\sigma'(z)} \underbrace{(a_j^L - t_j)}_{\frac{\partial C}{\partial a_j^L}}$$

$$\frac{\partial C}{\partial w_{jk}^L} = \delta_j^L a_k^{L-1}$$

$$\delta_j^L = \frac{\partial C}{\partial r_j^L} \frac{\partial r_j^L}{\partial z_j^L} = \frac{\partial C}{\partial r_j^L}$$

$$\delta_j^L = \Delta'(z_j^L) \frac{\partial C}{\partial a_j^L}$$

expressions for layer L
 what about layer L-1?

$$L \rightarrow l$$

$$\delta_j^l = \frac{\partial C}{\partial z_j^l} = \sum_k \frac{\partial C}{\partial z_k^{l+1}} \frac{\partial z_k^{l+1}}{\partial z_j^l}$$

$$= \sum_k \delta_k^{l+1} \frac{\partial z_k^{l+1}}{\partial z_j^l}$$

$$z_j^{l+1} = \sum_{i=1}^{N_l} w_{ij}^{l+1} a_i^l + b_j^{l+1}$$

$$\Rightarrow \frac{\partial z_k^{l+1}}{\partial z_j^l} = w_{kj}^{l+1} \sigma'(\bar{z}_j^l)$$

$$\delta_j^l = \sum_k \delta_k^{l+1} w_{kj}^{l+1} \Delta'(z_j^l)$$

updates-

$$w_{jk}^l \leftarrow w_{jk}^l - \eta \delta_j^l a_k^{l-1}$$

$$b_j^l \leftarrow b_j^l - \eta \frac{\partial C}{\partial b_j^l}$$

$$= b_j^l - \eta \delta_j^l$$

Algorithm

- Define architecture
- initialize $\{W, b\} = \Theta$
- Define nodes, hidden layers
+ activation functions
- hyperparameters
- learning rate + gradient
approx
- Define cost functions

- set up x
- Perform first feed forward
- continue to layer L
- $$a^L = \sigma^L(\sigma^{L-1}(\dots \sigma^1(z^1)\dots))$$
- compute δ_j^L
- Then Backpropagate

$$l = L-1, L-2, \dots, 1$$

$$\delta_j^l = \sum_k \delta_k^{l+1} w_{kj}^{l+1} \sigma'(z_j^l)$$

Training (Gradient part)

$$w_{jk}^l \leftarrow w_{jk}^l - \eta \delta_j^l a_k^{l-1}$$

$$b_j^l \leftarrow b_j^l - \eta \delta_j^l$$

Continue till convergence