

Neural Networks

The perceptron

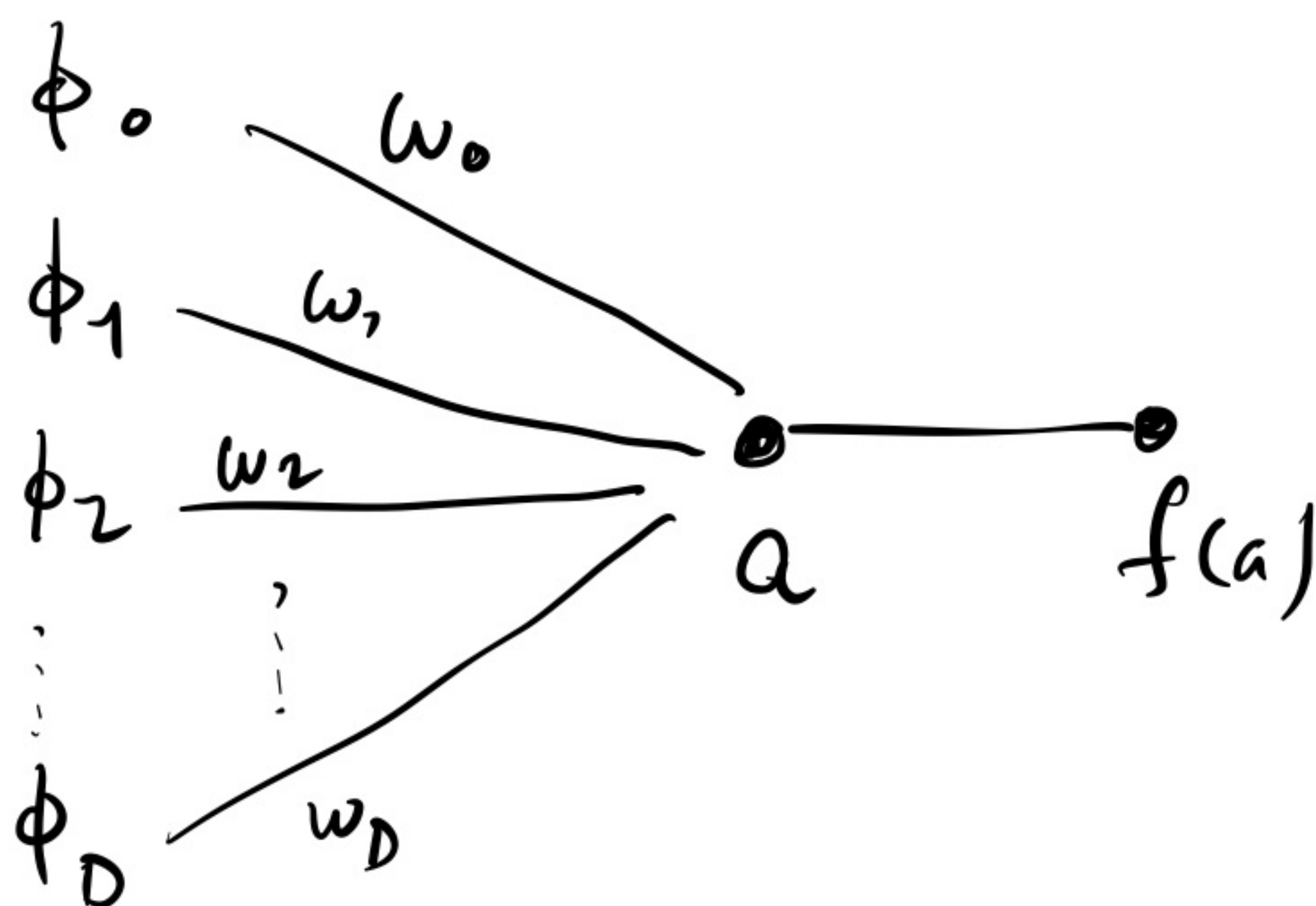
↳ originally → binary classification

$$y(\vec{x}) = f(\vec{w}^T \cdot \vec{\phi}(\vec{x})) = p(C_1 | \vec{x}, \vec{w})$$

↳

$$f(a) = \begin{cases} +1 & a \geq 0 \\ -1 & a < 0 \end{cases}$$

Convenient to label as $t_i = \{-1, 1\}$



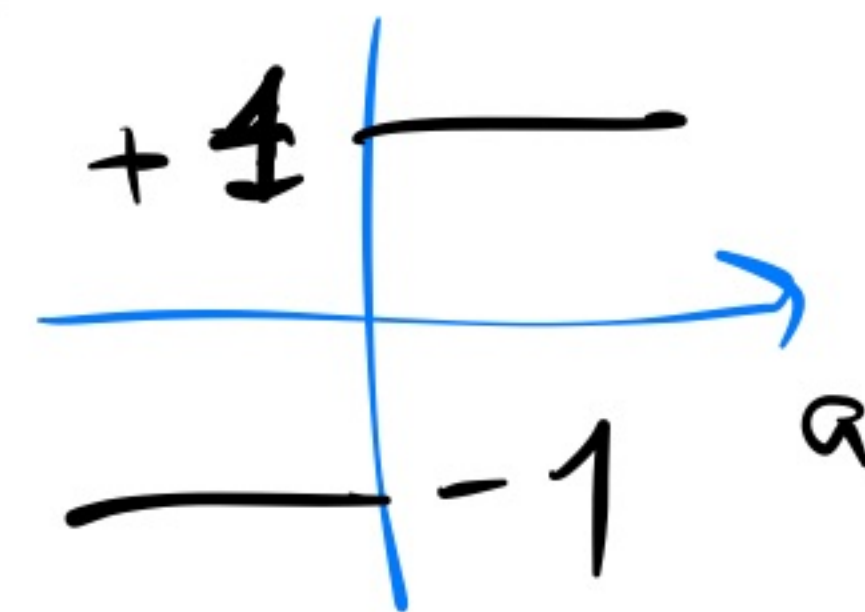
reminds of a
NN
(also
logistic
Regression)

To train the perceptron model

↳ minimise some cost function

↓ cross-entropy : problematic

$$\Rightarrow \frac{df}{da} \Rightarrow \text{discontinuous}$$



Instead :

$$E_P(\vec{w}) = - \sum_{i \in M} \vec{w}^T \cdot \vec{\phi} \cdot t_i$$

↳ misclassified points only

Note: if $t_i = +1 \Rightarrow$ a correct prediction has $\vec{w}^T \vec{\phi}_i > 0$

$t_i = -1 \Rightarrow$ " " "

$$\vec{w}^T \vec{\phi}_i < 0$$

↓ $\vec{w}^T \vec{\phi} t_i > 0 \quad \forall$ points correctly classified

misclassified points have $\vec{w}^T \vec{\phi}_i t_i < 0$

$E_P(\vec{w})$ is positive defined

How to minimise $E_p(\vec{w})$?

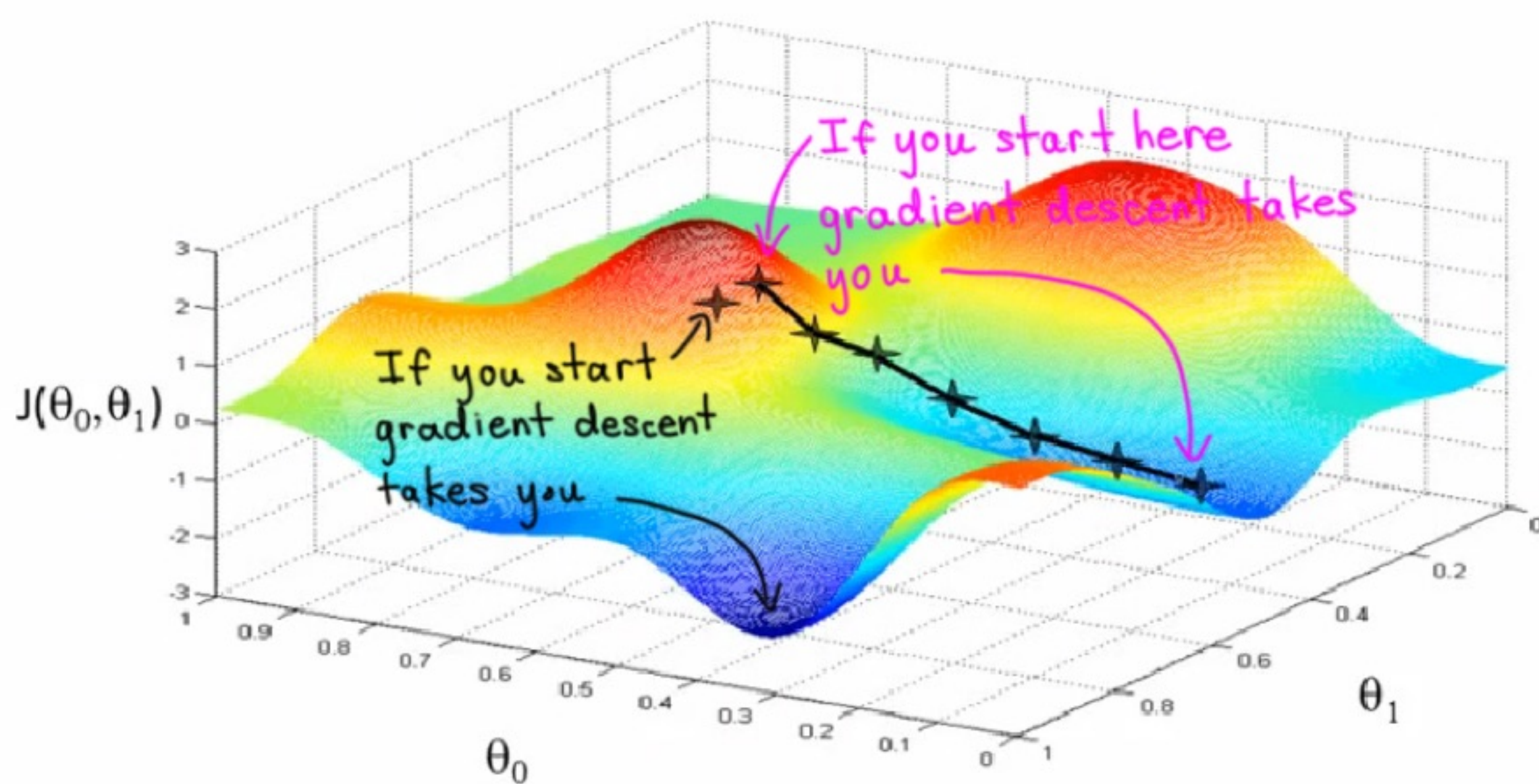
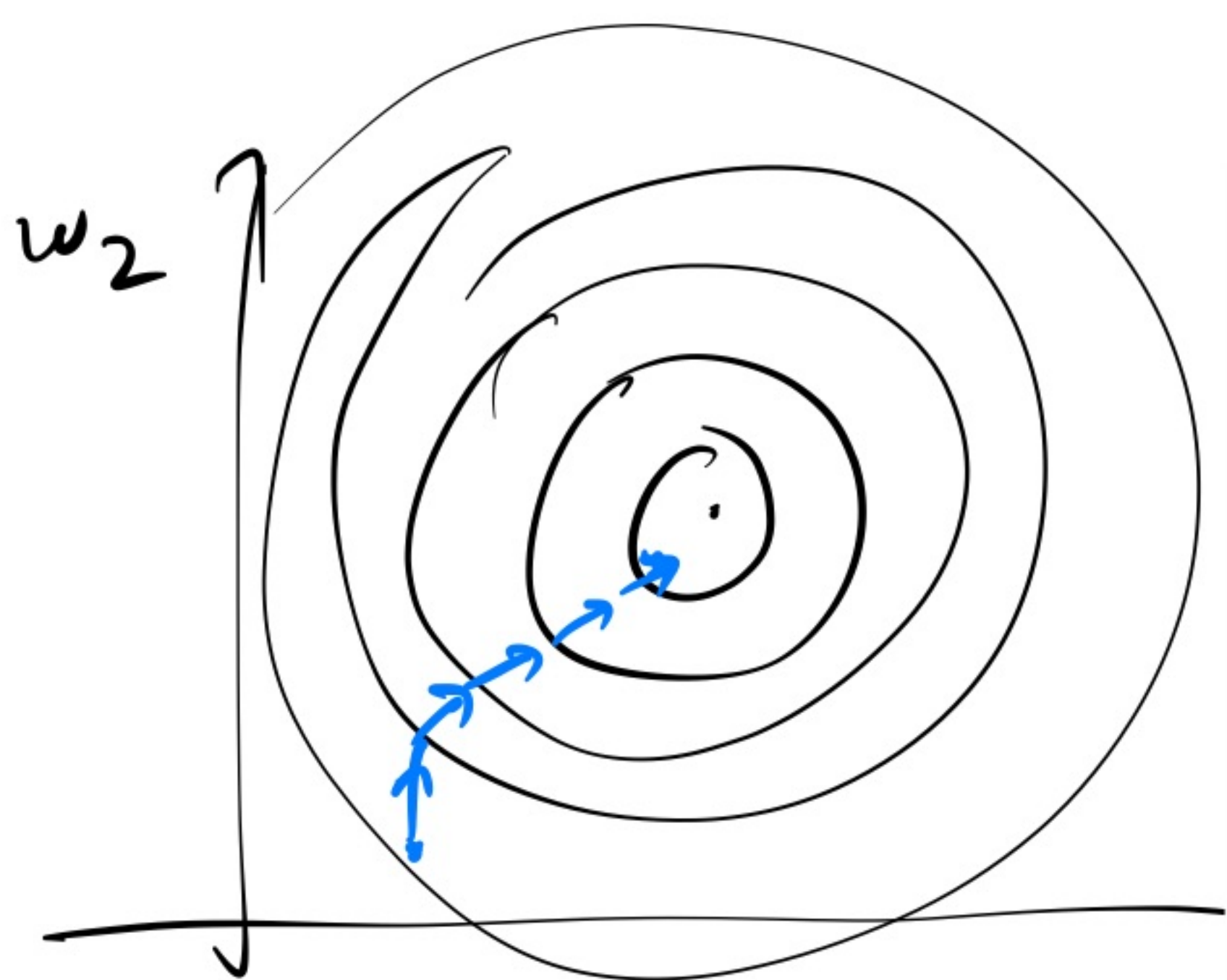
↳ Stochastic Gradient Descent

Gradient Descent

Iterative algo. for optimising a function
at iteration $\tau+1$ ("epoch")

$$\vec{w}_{\tau+1} = \vec{w}_{\tau} - \eta \frac{dC(\vec{w})}{d\vec{w}}$$

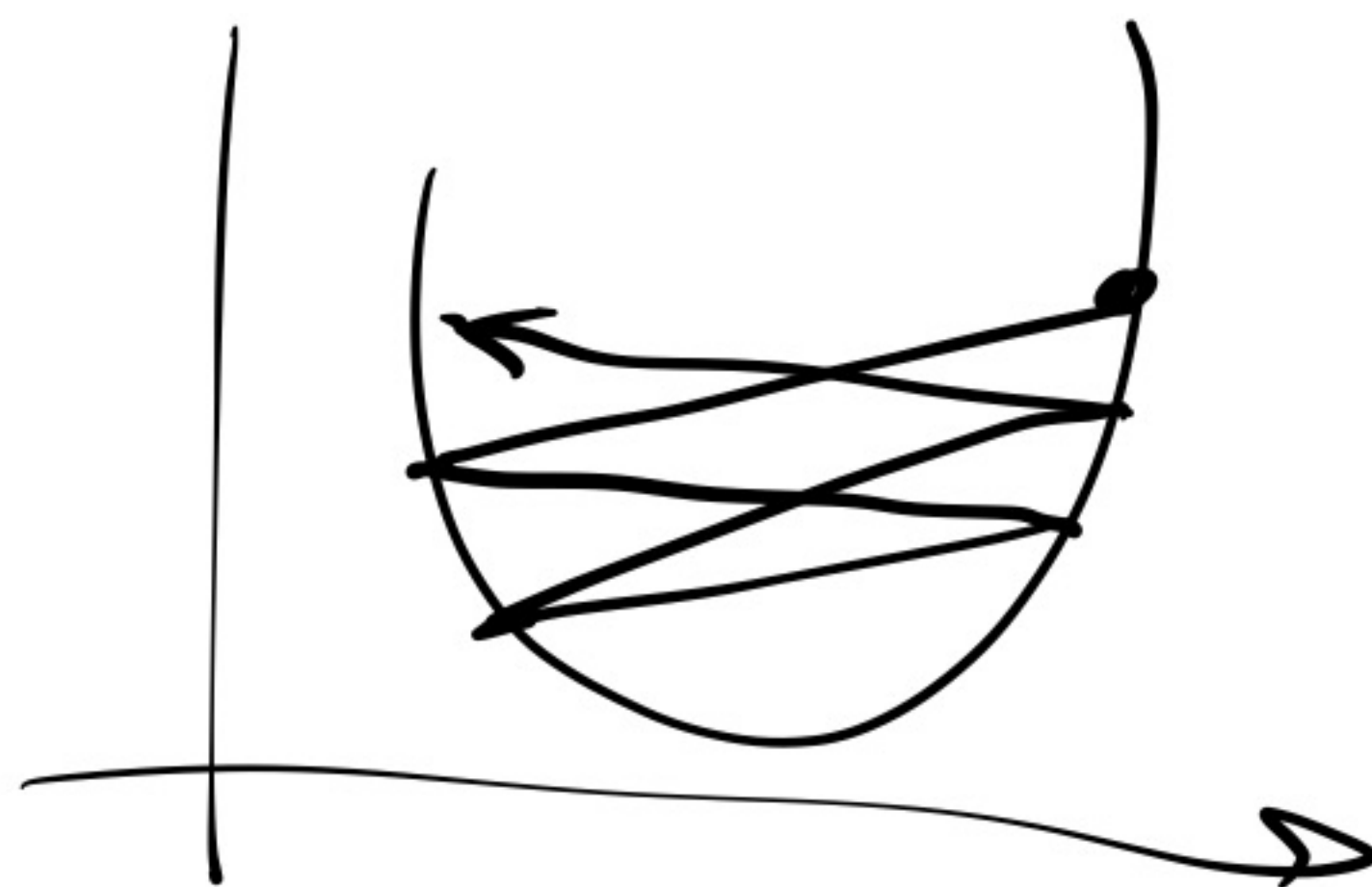
cost function
learning rate



↳ problem with local minima

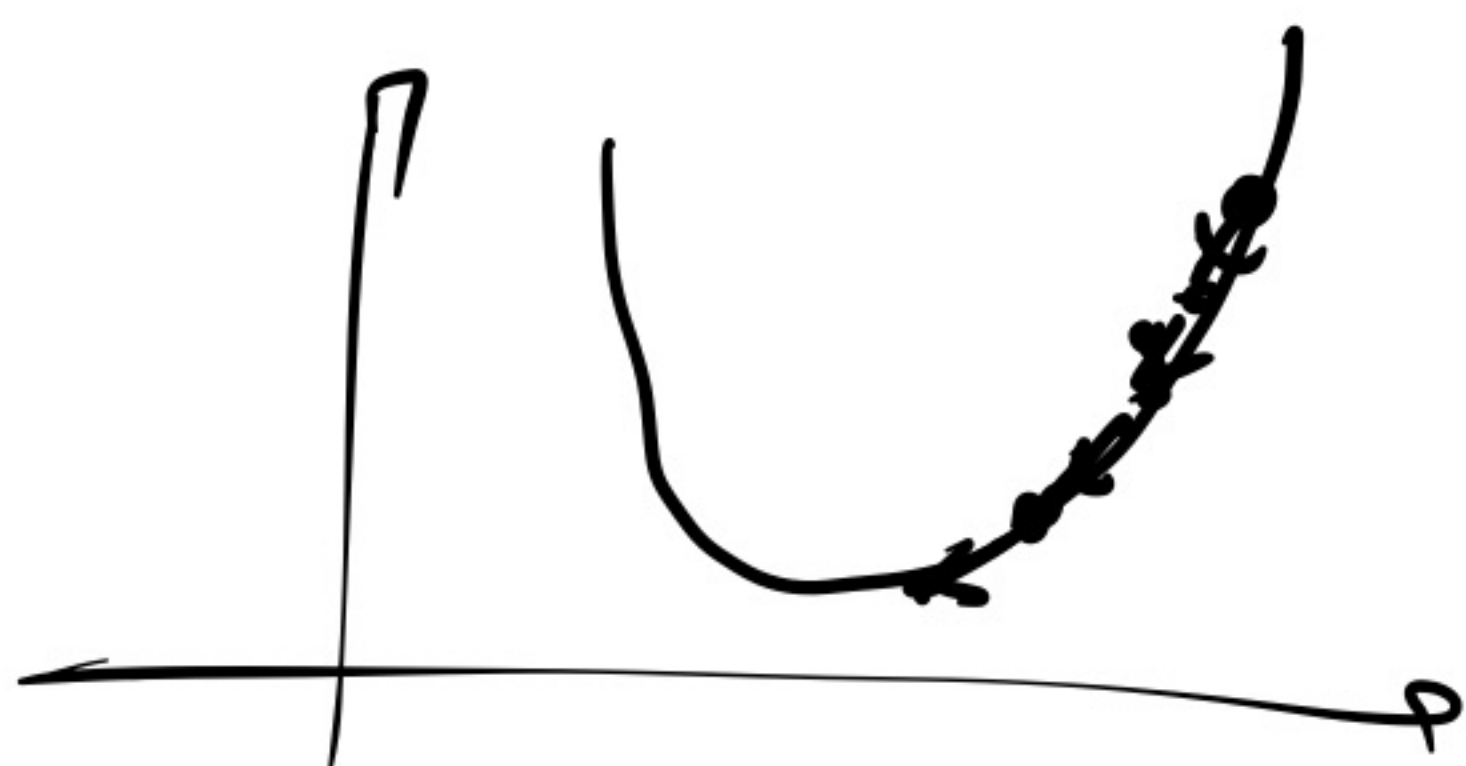
Depending on learning rate

large η



miss the minimum

too small



too long training time

η can be a constant $\forall$ epochs (τ)

but sometimes

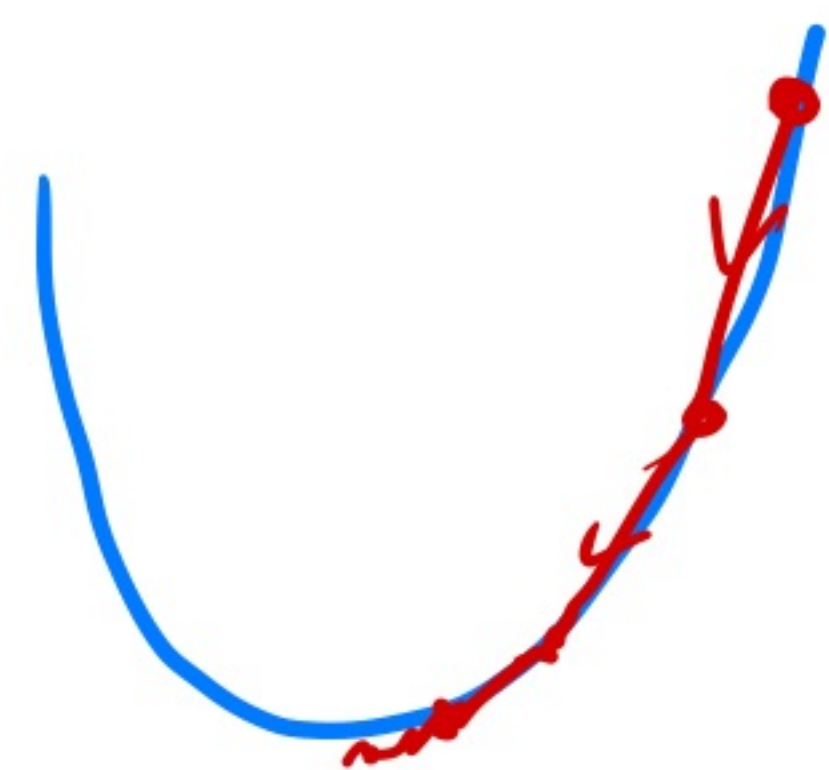
$$\eta \propto \frac{1}{\tau}$$

hyperparameter that should be optimised

Stochastic G.D:

approximate

$$\frac{dC(\vec{w})}{d\vec{w}}$$



$$C(\vec{w}) = \frac{1}{N} \sum_{i=1}^N C_i(\vec{w})$$

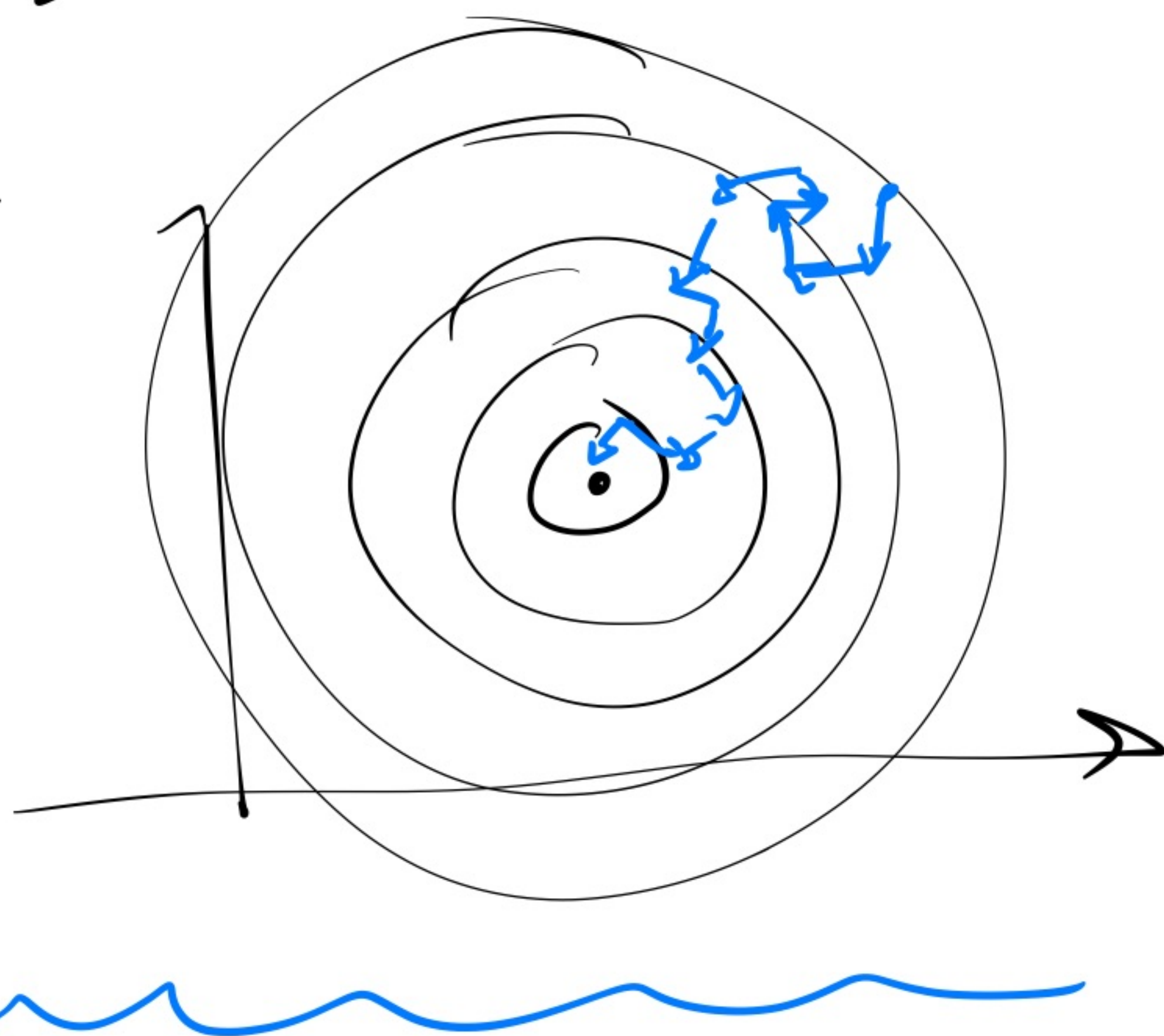
expensive for large N

$$\frac{dC(\vec{w})}{d\vec{w}} \approx \frac{dC_i(\vec{w})}{d\vec{w}}$$

chosen randomly

works because, on average

$$\frac{1}{N} \sum_i \frac{dc_i}{d\vec{w}} = \frac{dc}{d\vec{w}}$$



Back to perceptron

SGD

$$\vec{w}_{\tau+1} = \vec{w}_{\tau} + \eta \vec{\phi}_i t_i$$

$$\frac{\vec{w}_{\tau+1}}{\eta} = \frac{\vec{w}_{\tau}}{\eta} + \vec{\phi}_i t_i$$

$$\Rightarrow \vec{w}_{\tau+1} = \vec{w}_{\tau} + \vec{\phi}_i t_i \quad \left\{ \begin{array}{l} f(a) \text{ is} \\ \text{invariant} \\ \text{under rescaling} \\ \text{of } \vec{w} \end{array} \right.$$

At each iteration τ ,

find a random misclassified point,

$$\vec{w}_{\tau+1} = \begin{cases} \vec{w}_{\tau} + \vec{\phi}_i & \text{if } i \in C_1 \\ \vec{w}_{\tau} - \vec{\phi}_i & \text{if } i \in C_2 \end{cases}$$

in iteration $\tau+1$ the $E_p(\bar{w})$ will be smaller than in τ

$$-\bar{w}_{\tau+1}^T \phi_i t_i = -(\bar{w}_{\tau} + \bar{\phi}_i t_i)^T \bar{\phi}_i t_i$$

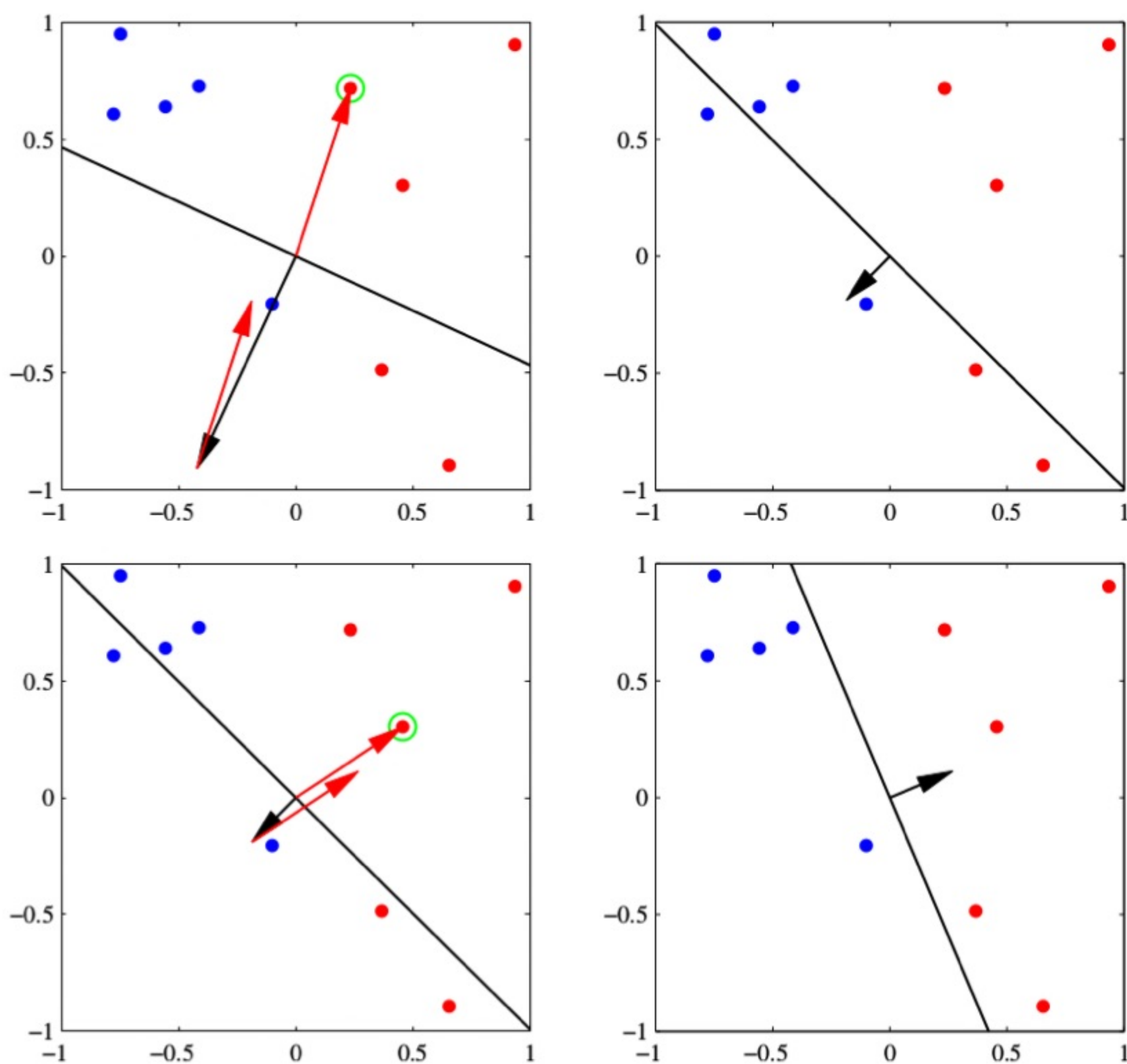
$$= -\bar{w}_{\tau}^T \bar{\phi}_i t_i$$

$$- \underbrace{(\bar{\phi}_i t_i)^T (\bar{\phi}_i t_i)}_{>0}$$

$$< -\bar{w}_{\tau}^T \bar{\phi}_i t_i$$

perceptron convergence theorem

data linearly separable



Neural nets

Start with linear model

$$y(\vec{x}, \vec{v}) = f\left(\sum_{h=1}^H v_h \phi_h(\vec{x})\right)$$

more
flexible
model

$f(a) = a$ for regression problems

$f(a)$: non-linear function
e.g. sigmoid

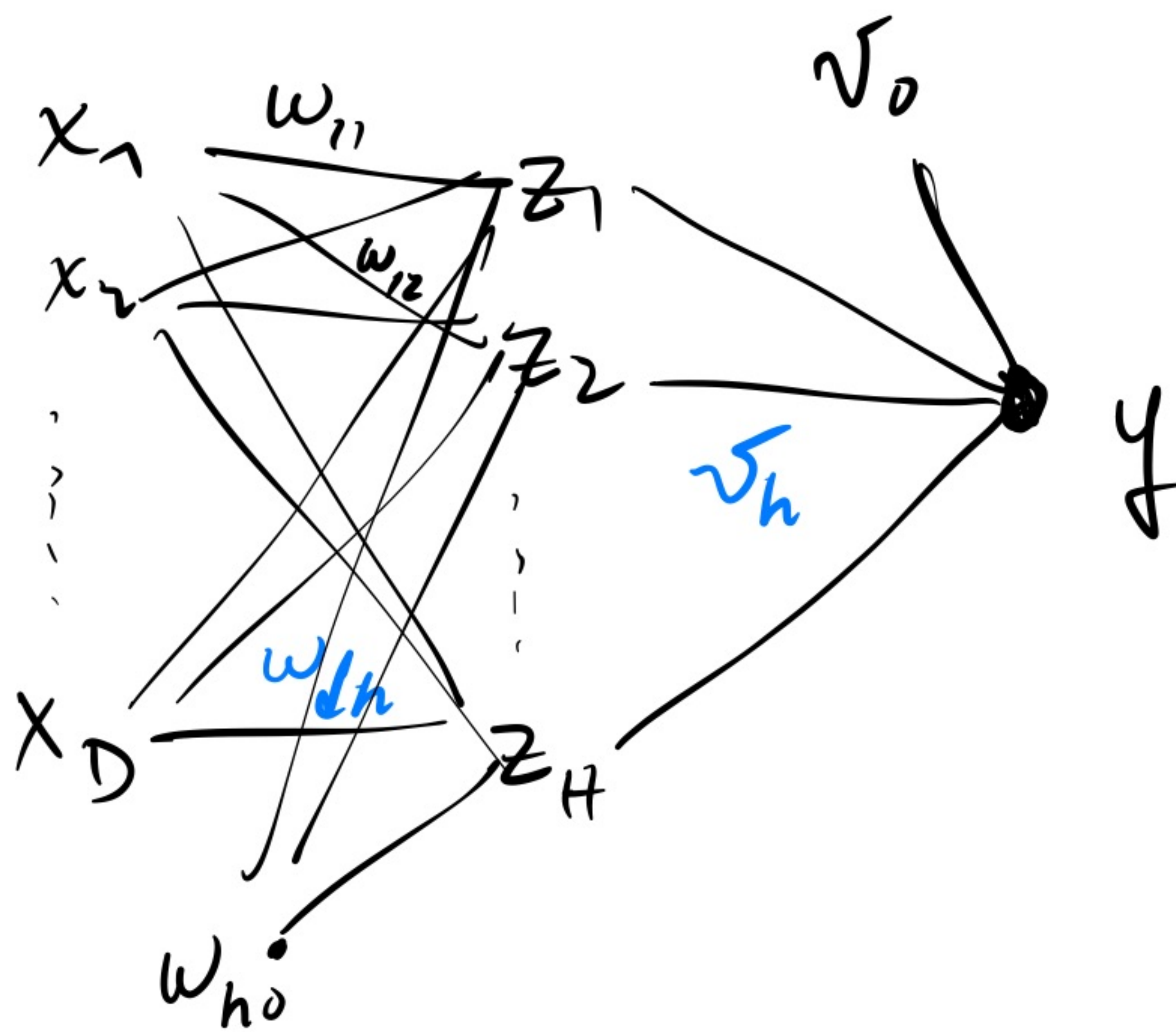
$$p(C_1 | \vec{x}, \vec{w})$$

$$\phi_h(\vec{x}, \vec{w}_h) = g\left(\sum_{d=1}^D w_{hd} x_d + w_{h0}\right)$$

non-linear (activation)
function

$$y(\vec{x}, \vec{v}, \vec{w}) = f\left(\sum_{h=1}^H v_h \underbrace{g\left(\sum_{d=1}^D w_{hd} x_d + w_{h0}\right)}_{\equiv z_h} + v_0\right)$$

one-hidden layer neural network
(feedforward)



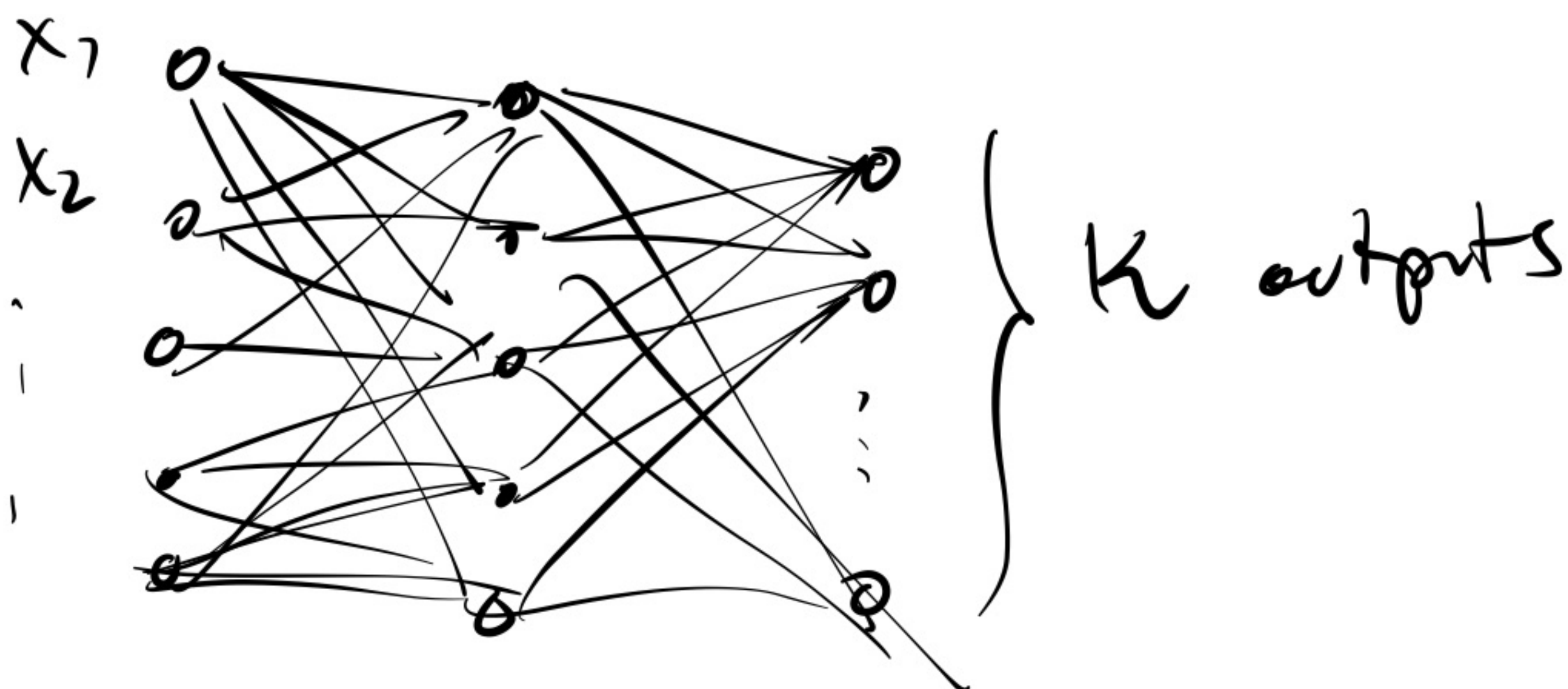
For classification

(output is k dimensional)

$$\vec{t}_i = \{t_1, \dots, t_k\}$$

$$k=1, \dots, k$$

$$y_k(\vec{x}, \vec{v}, \vec{w}) = f\left(\sum_{h=1}^{H+1} v_{hk} g\left(\sum_d w_{hd} x_d + w_{h0}\right) + v_{k0}\right)$$



$$y_k(\vec{x}, \vec{v}, \vec{w}) = f\left(\sum_{h=1}^{H+1} v_{hk} g\left(\sum_d w_{hd} x_d + w_{h0}\right) + v_{k0}\right)$$

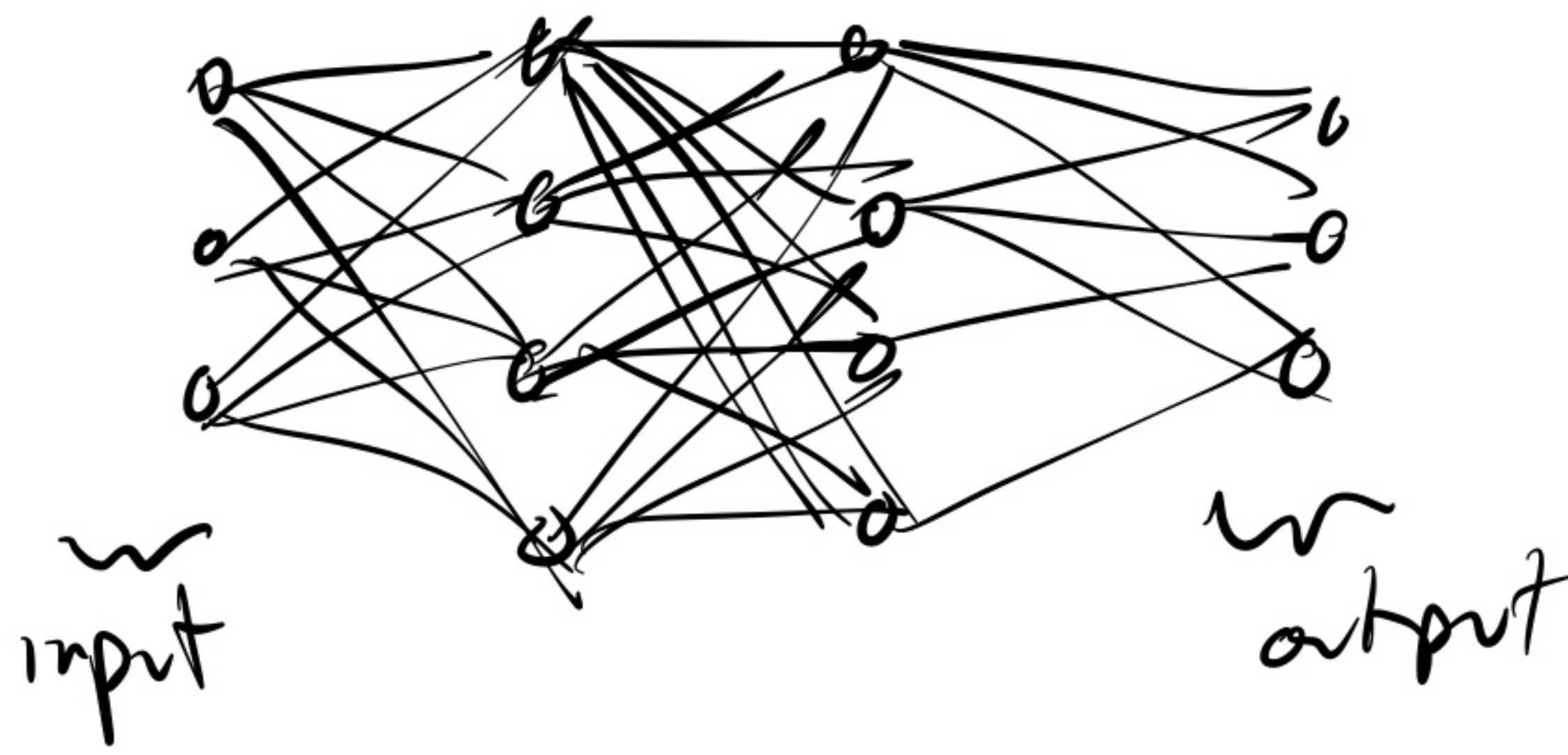
$$Y: N \times K, \quad V = H \times K, \quad V_0 = N \times K$$

$$X: N \times D \quad W: D \times H \quad W_0: N \times H$$

$$Y = f\left(\underbrace{g\left(\underbrace{XW}_{N \times H} + \underbrace{W_0}_{N \times H}\right)}_{(N \times H)} \underbrace{V}_{H \times K} + \underbrace{V_0}_{N \times K}\right)$$

$N \times K$

A NN with 2 hidden layers



$$Y = f \left(g_2 \left[g_1 (XW^{(1)} + W_0^{(1)}) W^{(2)} + W_0^{(2)} \right] V + V_0 \right)$$

Activation functions

Sigmoid $\Rightarrow$ hidden layers

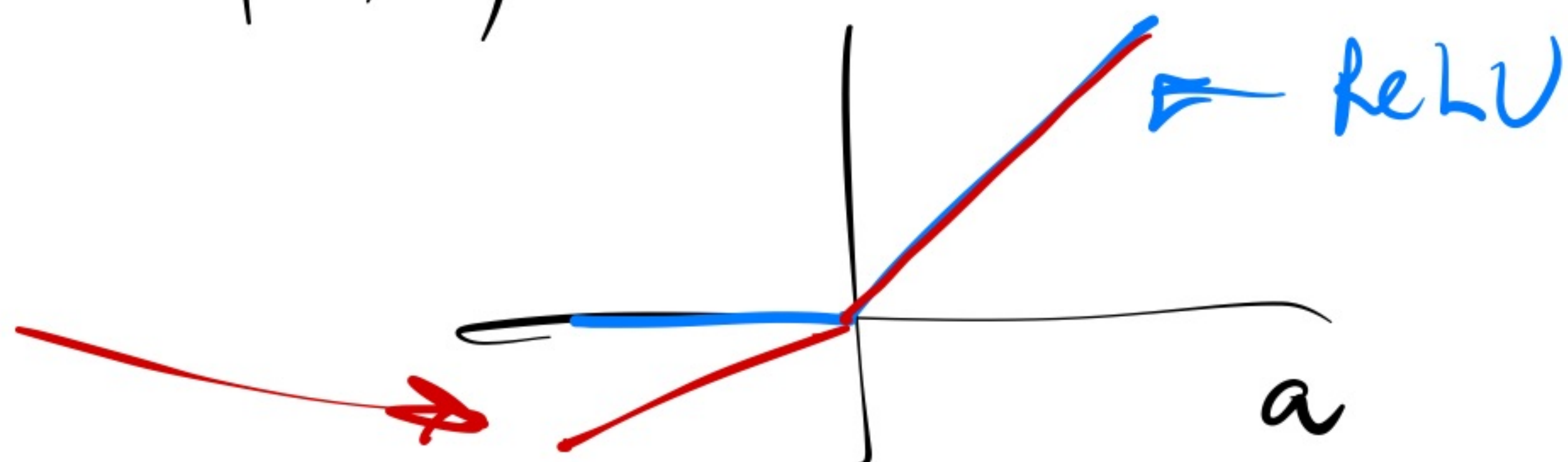
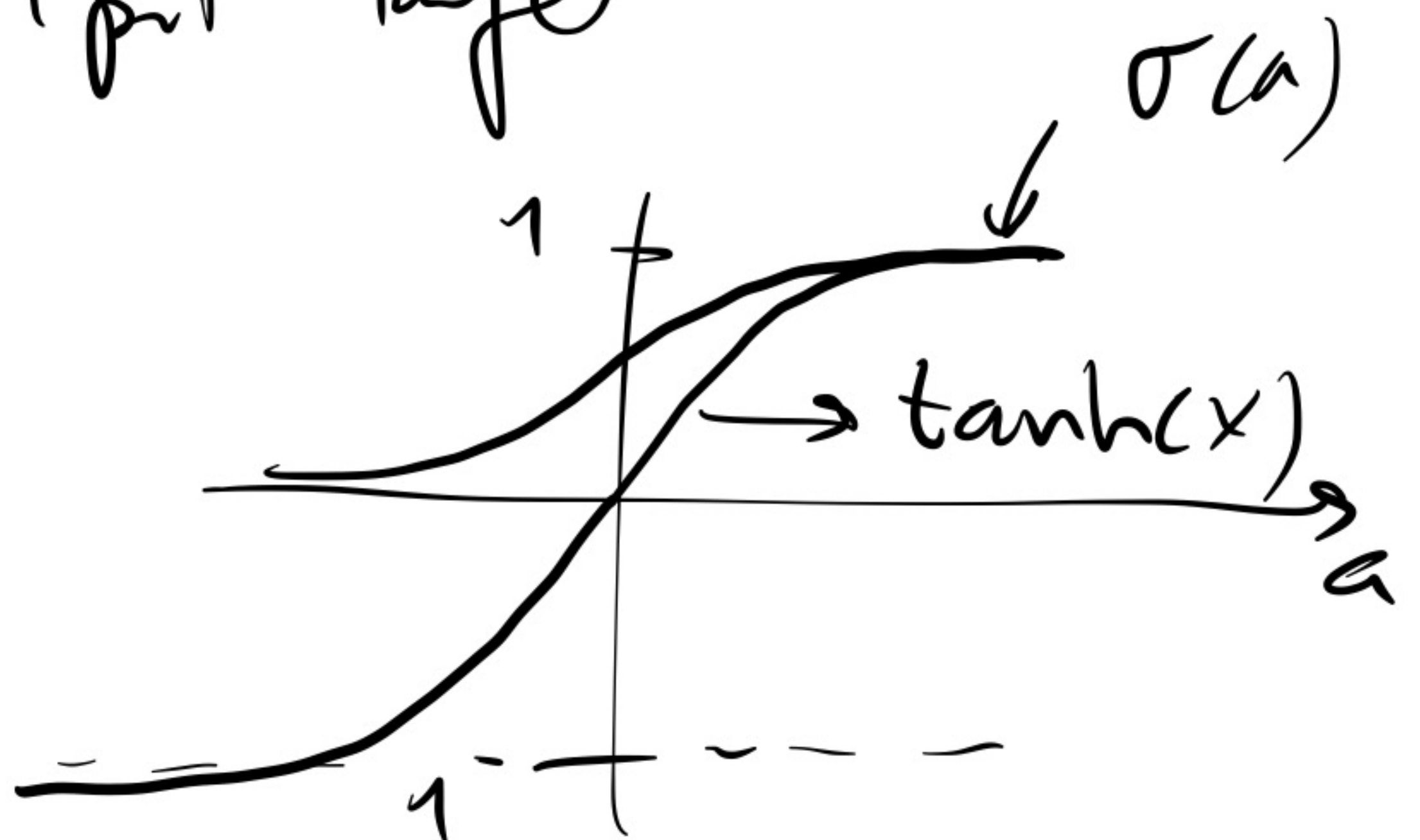
softmax $\Rightarrow$ output layer

tanh, ...

more recently

$$\text{ReLU} = \max(0, a)$$

leaky ReLU



All the activation function give in general
good results



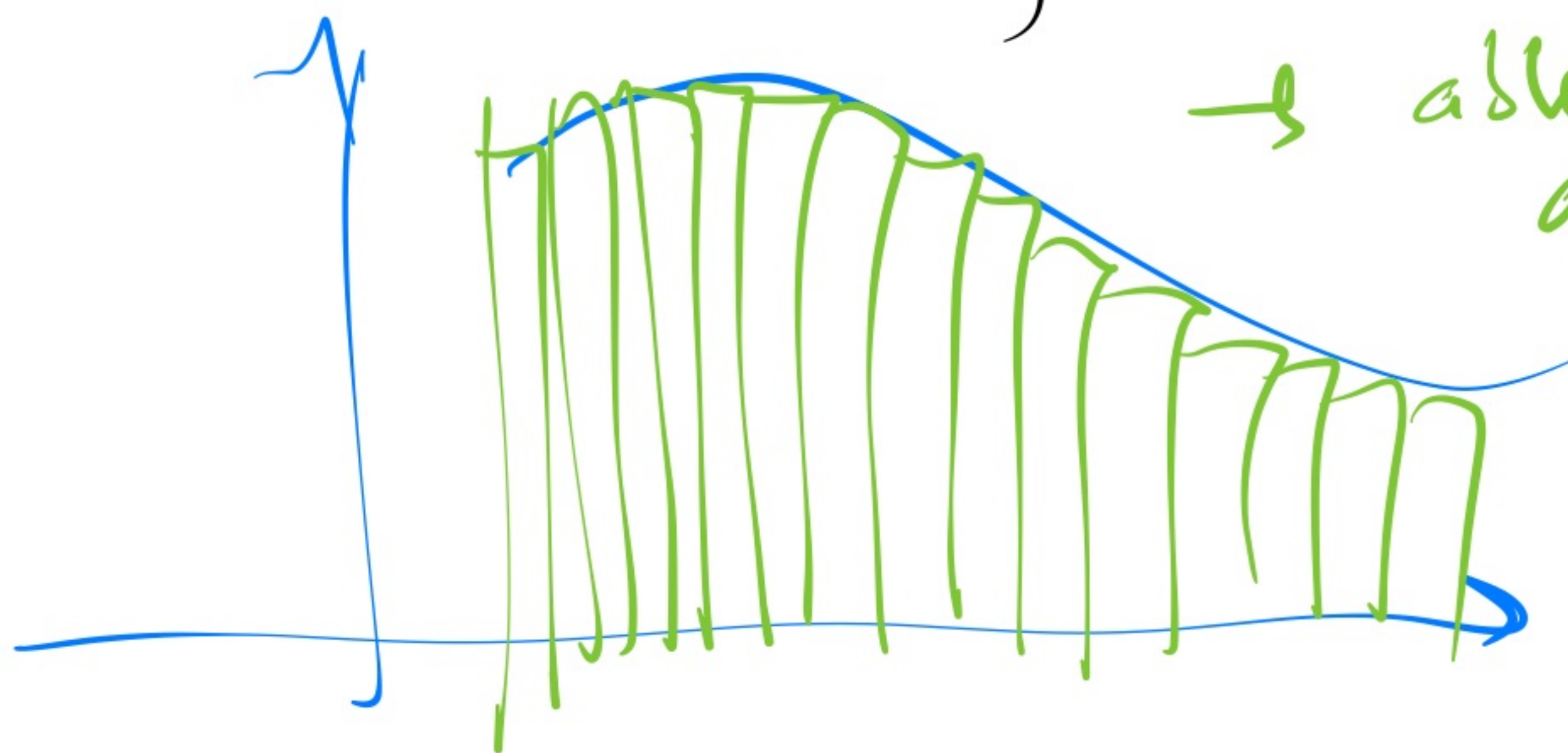
The Universal Approximation Theorem

A feedforward NN with a single hidden layer with finite # of units can approximate any continuous function under mild assumptions of the activation function

originally
Proven with

$$\sigma(a) = \begin{cases} 1 & a \rightarrow \infty \\ 0 & a \rightarrow -\infty \end{cases}$$

but
e.g. ReLU works equally well



→ able to
approx with
rectangles
↓
good

Training the network

→ find optimum weights that
maximise likelihood (minimise cost
function)

↓
Classification

$$\mathcal{E}(\vec{w}) = - \sum_{i=1}^N \sum_{k=1}^K t_{ik} \ln \underbrace{y_k(\vec{x}_i, \vec{w})}_{\equiv y_{ik}}$$

softmax sigmoid

$$Y = \sigma_Y \left(\sigma(XW + W_0)V + V_0 \right)$$

$$\equiv \sigma_Y(ZV + V_0)$$

sigmoid

$$\sigma(a) = \frac{1}{1 + e^{-a}} \Rightarrow \frac{d\sigma}{da} = \sigma(1 - \sigma)$$

softmax

$$\sigma(a_k) = \frac{e^{a_k}}{\sum_j e^{a_j}} \Rightarrow \frac{d\sigma(a_k)}{da_j} = \sigma_k(1 - \sigma_k) \delta_{kj} \\ = \sigma_k(\delta_{kj} - \sigma_j)$$

$$\frac{\partial \mathcal{E}}{\partial v_{ok}} = - \sum_i \sum_{k'} t_{ik'} \frac{1}{y_{ik'}} \frac{\partial y_{ik'}}{\partial v_{ok}}$$

$$= - \sum_i \sum_{k'} t_{ik'} \frac{1}{\cancel{y_{ik'}}} \cancel{y_{ik'}} (\delta_{kk'} - y_{ik}) \frac{\partial a_{ik}}{\partial v_{ok}}$$

1

$$= - \sum_i \sum_{k'} t_{ik'} (\delta_{kk'} - y_{ik})$$

$$= - \sum_i \sum_{k'} (t_{ik'} y_{ik}) - \sum_i t_{ik}$$

$$= - \sum_i y_{ik} \underbrace{\sum_{k'} t_{ik}}_{=1} - \sum_i t_{ik}$$

$$= \sum_i (y_{ik} - t_{ik}) = \sum_i (Y - T)_{ik}$$

Next

$$\frac{\partial \mathcal{E}}{\partial v_{hk}} = - \sum_i \sum_{k'} t_{ik'} (\delta_{kk'} - y_{ik}) \frac{\partial a_{ik}}{\partial v_{hk}}$$

$$= - \sum_i (y_{ik} - t_{ik}) z_{ih}$$

z_{ih}

$$\frac{\partial \mathcal{E}}{\partial V} = Z^T (Y - T)$$

$$\frac{\partial \mathcal{E}}{\partial w_{ok}} = - \sum_i \sum_k t_{ik} \frac{1}{y_{ik}} (1 - y_{ik}) \frac{\partial a_{ik}}{\partial z_{ih}} \frac{\partial z_{ih}}{\partial w_{ok}}$$

$$= \sum_i \sum_k (t_{ik} - y_{ik}) v_{hk} z_{ih} (1 - z_{ih})$$

$$= - \sum_i \sum_k (T - Y)_{ik} (V^T)_{kh} z_{ih} (1 - z_{ih})$$

$$= - \sum_i \left[(Y - T) V^T * Z * (1 - Z) \right]_{ih}$$

↑
Schur product
(element-by-element)

finally

$$\frac{\partial \mathcal{E}}{\partial w_{oh}} = \sum_i \left[(Y - T) V^T * Z * (1 - Z) \right]_{ih} x_{id}$$

$$= X^T \left[(Y - T) V^T * Z * (1 - Z) \right]$$