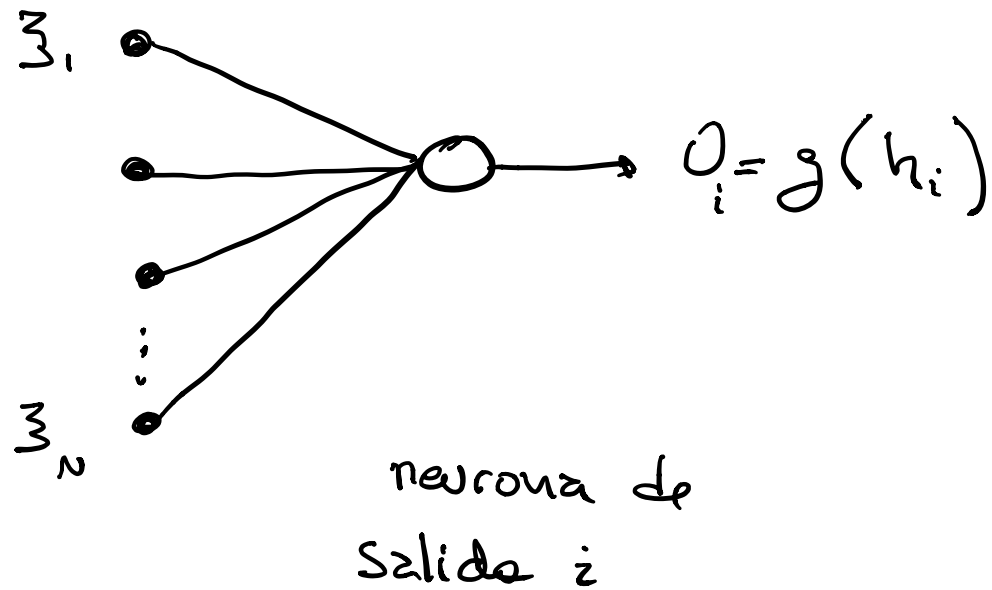


Seguimos con perceptron simple



Ayer vimos el caso $g(h) = \text{signo}(h)$

Si tenemos un conjunto de entrenamientos de P relaciones input-output correctas

$$\vec{x}_i^{\mu} \rightarrow d_i^{\mu} \quad \mu = 1, 2, \dots, P$$

Si el problema es linealmente separable, tenemos un algoritmo para $g(h) = \text{signo}(h)$

$$w_{ik}^{\text{new}} = w_{ik}^{\text{old}} + \eta (d_i^M - O_i^M) \sum_k^M$$

Neuronas de salida lineales

$$g(x) = x$$

$$O_i = \sum_k w_{ik} \sum_k$$

Otro vez queremos que
tenemos un conjunto de
entrenamiento

$$\sum^M \rightarrow d_i^M$$

$$d_i^M = \sum_k w_{ik} \sum_k^M$$

Dado un vector cualquiera $\bar{w}$
puedo calcular el error que
comete al aprender

$$E(\bar{w}) = \frac{1}{2} \sum_{i, \mu} (d_i^\mu - O_i^\mu)^2$$

$$= \frac{1}{2} \sum_{\mu} \sum_i (d_i^\mu - O_i^\mu)^2$$

$$E(\bar{w}) = \frac{1}{2} \sum_{\mu} \sum_i \left(d_i^\mu - \sum_k w_{ik} z_k^\mu \right)^2$$

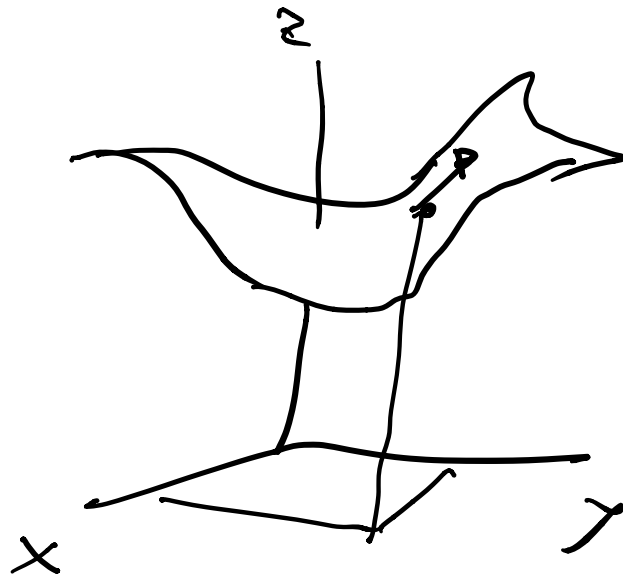
Agregué muchas neuronas
de salida etiquetadas con i

Gradiente

$$\text{Sea } f: \mathbb{R}^2 \rightarrow \mathbb{R}$$

$$f(x_1, x_2) = x_1 + \sin(x_2)$$

$$\nabla f_{\bar{x}} = \left(\frac{\partial f}{\partial x_1}, \frac{\partial f}{\partial x_2} \right)_{\bar{x}}$$



El gradiente define la
dirección en un punto en $\mathbb{R}^2$
en la cual f crece más.

f crece más rápido en

la dirección $-\vec{\nabla} f$

En el más general

$$f(x_1, x_2, \dots, x_n)$$

En nuestro caso

$$E(\tilde{w}) = E(w_{11}, w_{12}, \dots, w_{21}, w_{22}, \dots, w_{nm})$$

m : no de neuronas de salida

$$\nabla E_{ik} = \frac{\partial E}{\partial w_{ik}}$$

$$w_{ik}^{\text{new}} = w_{ik}^{\text{old}} + \Delta w_{ik}$$

$$\Delta w_{ik} = -\eta \frac{\partial E(\bar{w})}{\partial w_{ik}}$$

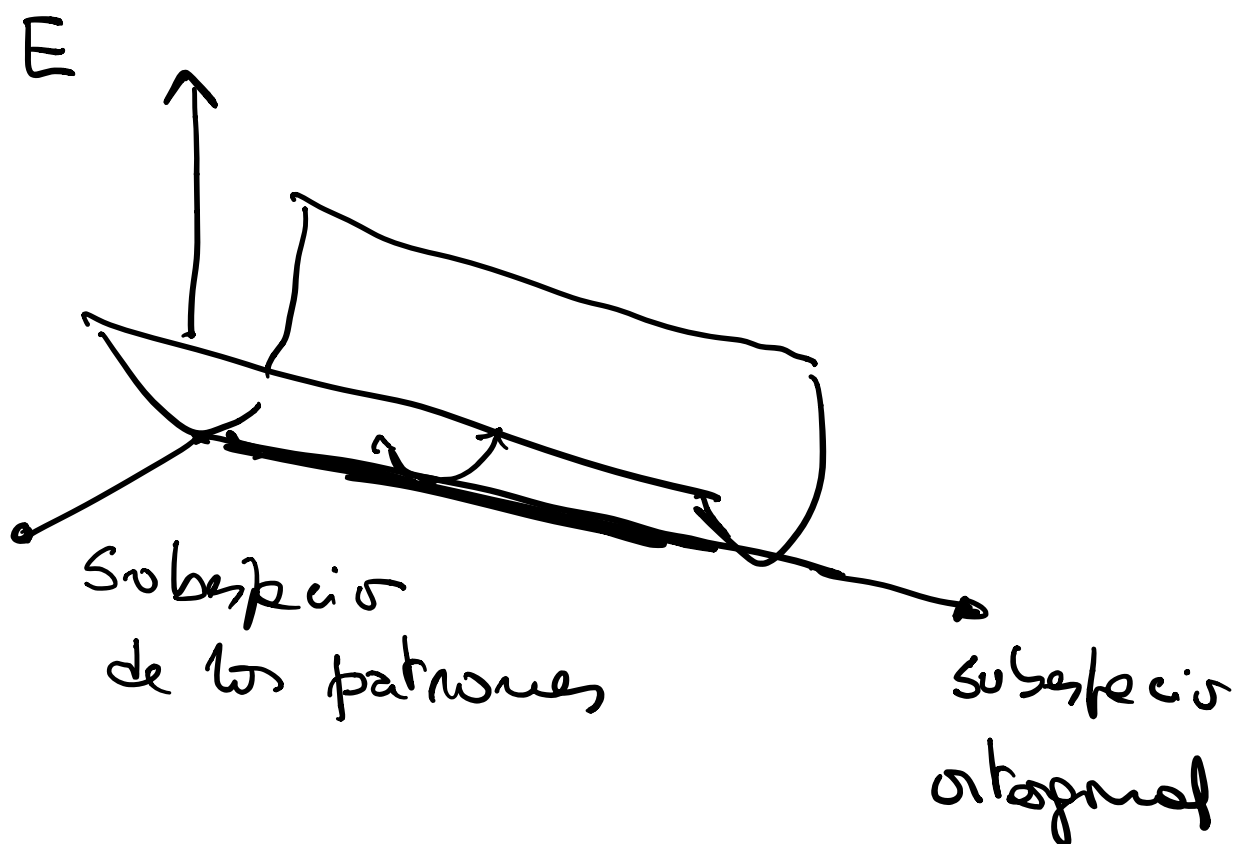
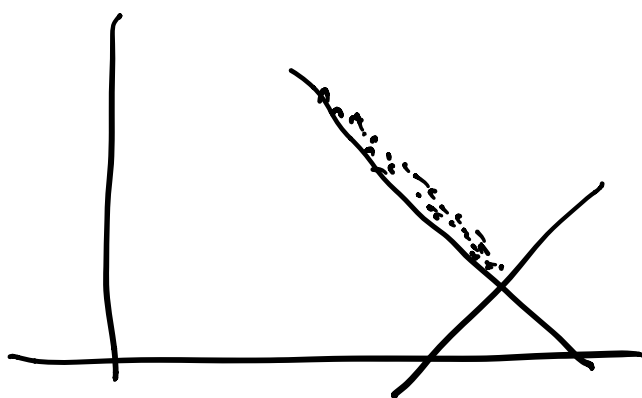
$$\Delta w_{ik} = +\eta \sum_{\mu} \underbrace{(d_i^{\mu} - \sum_k w_{ik} z_k^{\mu})}_{\text{error}} z_k^{\mu}$$

$$= \eta \sum_{\mu} (d_i^{\mu} - o_i^{\mu}) z_k^{\mu}$$

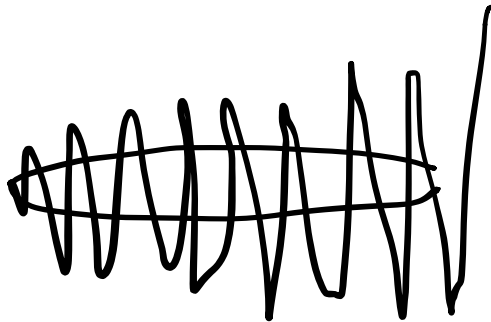
$$\Delta w_{ik} = \eta \sum_{\mu} \delta_i^{\mu} z_k^{\mu}$$

Se la conosce come $\delta_i^{\mu} = (d_i^{\mu} - o_i^{\mu})$

- regola delta
- regola xdelive
- Widrow-Hoff
- regola LMS



En general los mínimos no son puntos en $\mathbb{R}^n$ sino subespacios



El libro prueba la convergencia

— • —

Unidades de salida no lineales

$g(h)$ no lineal

- $g(h) = \frac{1}{1 + e^{-\beta h}}$

- $g(h) = \tanh(\beta h)$

la topología es la misma

Tenemos 3 ejemplos correctos

$$E(\bar{w}) = \frac{1}{2} \sum_{\mu} \sum_i (d_i^{\mu} - O_i^{\mu})^2$$

$$= \frac{1}{2} \sum_{\mu} \sum_i \left[d_i^{\mu} - g\left(\sum_k w_{ik} z_k^{\mu}\right) \right]^2$$

$$w_{ik}^{\text{new}} = w_{ik}^{\text{old}} + \Delta w_{ik}$$

$$\Delta w_{ik} = -\eta \frac{\partial E(\bar{w})}{\partial w_{ik}}$$

$$\frac{\partial E(\bar{w})}{\partial w_{ik}} = -\frac{1}{2} \sum_{\mu} 2 [d_i^{\mu} - o_i^{\mu}] g'(h_i^{\mu}) \sum_k^{\mu}$$

↑

$$\Delta w_{ik} = -\eta \frac{\partial E(\bar{w})}{\partial w_{ik}}$$

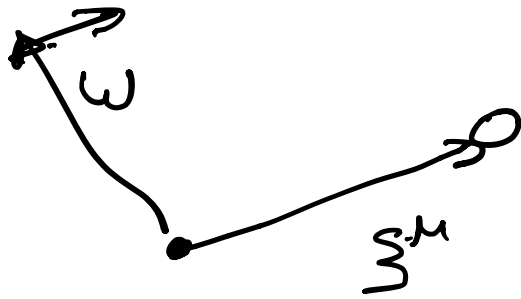
$$\Delta w_{ik} = \eta \sum_{\mu} \overbrace{[d_i^{\mu} - o_i^{\mu}] g'(h_i^{\mu}) \sum_k^{\mu}}$$

$$o_i^{\mu} = g\left(\sum_k w_{ik} \sum_k^{\mu}\right)$$

$$\Delta w_{ik} = \eta \sum_{\mu} \delta_i^{\mu} \sum_k^{\mu}$$

↑

Δw measured
as ξ^μ



$$\Delta \vec{w}_i = \eta \sum_{\mu} \delta_i^{\mu} \vec{\xi}^{\mu}$$

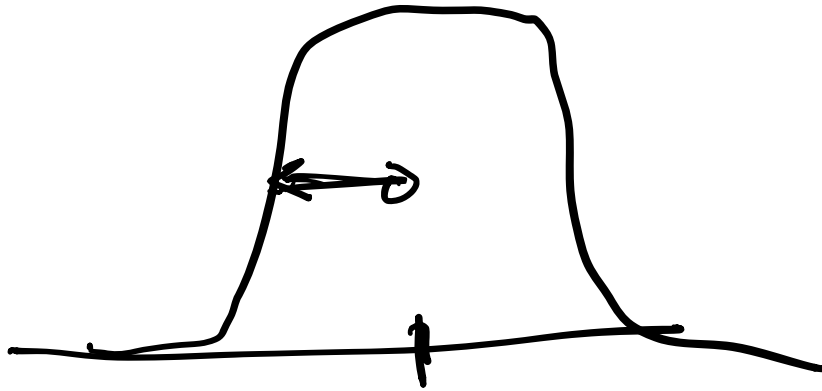
2: $g = \tanh(\beta h)$

$$g' = \beta(1 - g^2)$$

2: $g = \frac{1}{1 + e^{-2\beta h}}$

$$g' = 2\beta g(1 - g)$$



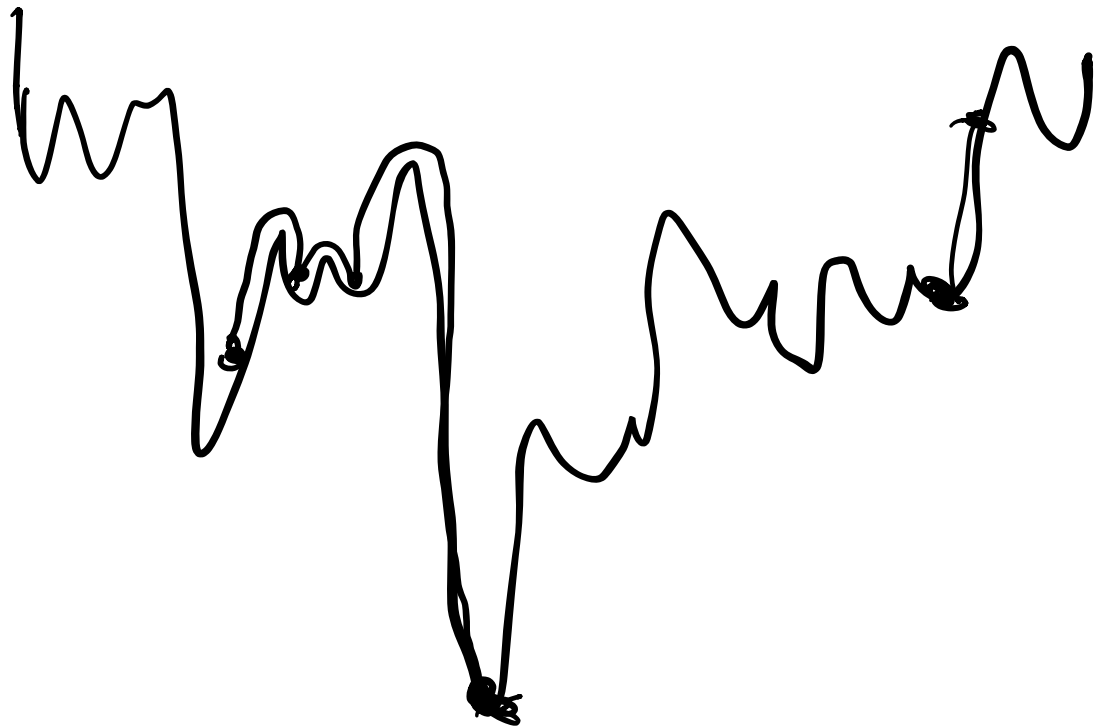


Una medida mejor del error
de la función

$$E(\bar{\omega}) = \sum_{i, \mu} \left[\frac{1}{2} (1 + d_i^\mu) \log \left(\frac{1 + d_i^\mu}{1 + o_i^\mu} \right) + \frac{1}{2} (1 - d_i^\mu) \log \left(\frac{1 - d_i^\mu}{1 - o_i^\mu} \right) \right]$$

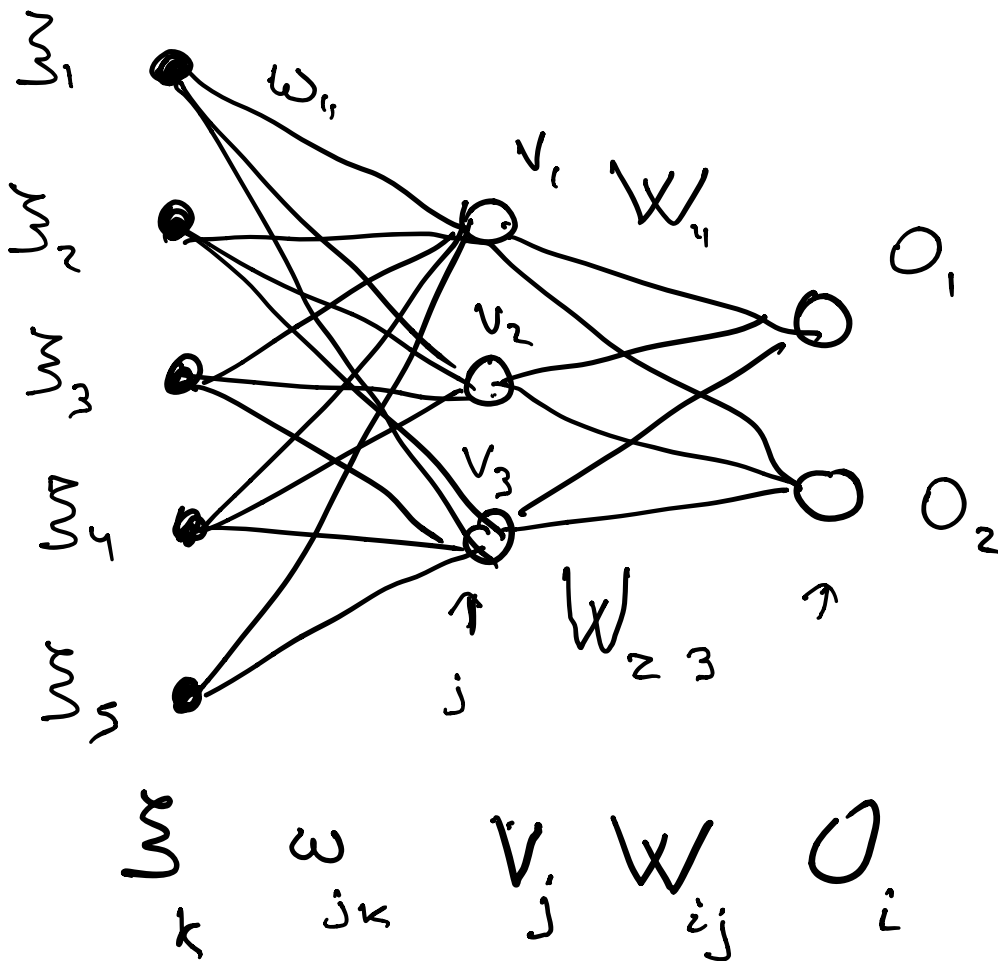
Es bueno ser neuronas estóricas

Coluceno el h y fabricano
una probabilidad $\neq$



Perceptron multi capas

Suponemos una capa oculta



Tenemos p ejemplos correctos

Presentamos un ejemplo μ

Calculamos todos los V_j

$$h_j^\mu = \sum_k w_{jk} z_k^\mu$$

$$V_j^\mu = g(h_j^\mu) = g\left(\sum_k w_{jk} z_k^\mu\right)$$

Calculamos los nodos

$$h_i^\mu = \sum_j w_{ij} V_j$$

$$= \sum_j w_{ij} g\left(\sum_k w_{jk} z_k^\mu\right)$$

le solido i' resé

$$O_i^\mu = g(h_i^\mu) = g\left(\sum_j w_{ij} V_j\right)$$

$$O_i^{\mu} = g \left(\sum_j W_{ij} g \left(\sum_k \omega_{jk} \xi_k^{\mu} \right) \right) \quad \leftarrow$$

Decí supuestos que los 2
capas tienen función g

Definimos una función error

$$E(\bar{w}) = \frac{1}{2} \sum_{\mu} \sum_i (d_i^{\mu} - O_i^{\mu})^2$$

$$= \frac{1}{2} \sum_{\mu} \sum_i \left(d_i^{\mu} - g \left(\sum_j W_{ij} g \left(\sum_k \omega_{jk} \xi_k^{\mu} \right) \right) \right)^2$$

$$\Delta W_{ij} = -\eta \frac{\partial E}{\partial W_{ij}}$$

$$= \eta \sum_{\mu} [d_i^{\mu} - O_i^{\mu}] g'(h_i^{\mu}) V_j^{\mu}$$

$$= \eta \sum_{\mu} \delta_i^{\mu} V_j^{\mu}$$

$$\delta_i^{\mu} = [d_i^{\mu} - O_i^{\mu}] g'(h_i^{\mu})$$

Vamos a la otra capa de
acoplamiento

$$\Delta \omega_{jk} = - \eta \frac{\partial E}{\partial \omega_{jk}}$$

$$= - \eta \sum_{\mu} \frac{\partial E}{\partial V_j} \cdot \frac{\partial V_j}{\partial \omega_{jk}}$$

$$= \eta \sum_i (d_i^{\mu} - O_i^{\mu}) g'(h_i^{\mu}) \times$$

$$\times W_{ij} g'(h_j^{\mu}) \sum_k^{\mu}$$

$$= \eta \sum_{\mu} \sum_i \delta_i^{\mu} W_{ij} g'(h_j^{\mu}) \sum_k \delta_k^{\mu}$$

$$= \eta \sum_{\mu} \delta_j^{\mu} \sum_k \delta_k^{\mu}$$

donc $\delta_j^{\mu} = g'(h_j^{\mu}) \sum_i W_{ij} \delta_i^{\mu}$

le information va de entrée
à sortie

le connexion va de la sortie à
la entrée

BACK PROPAGATION

Ejemplo

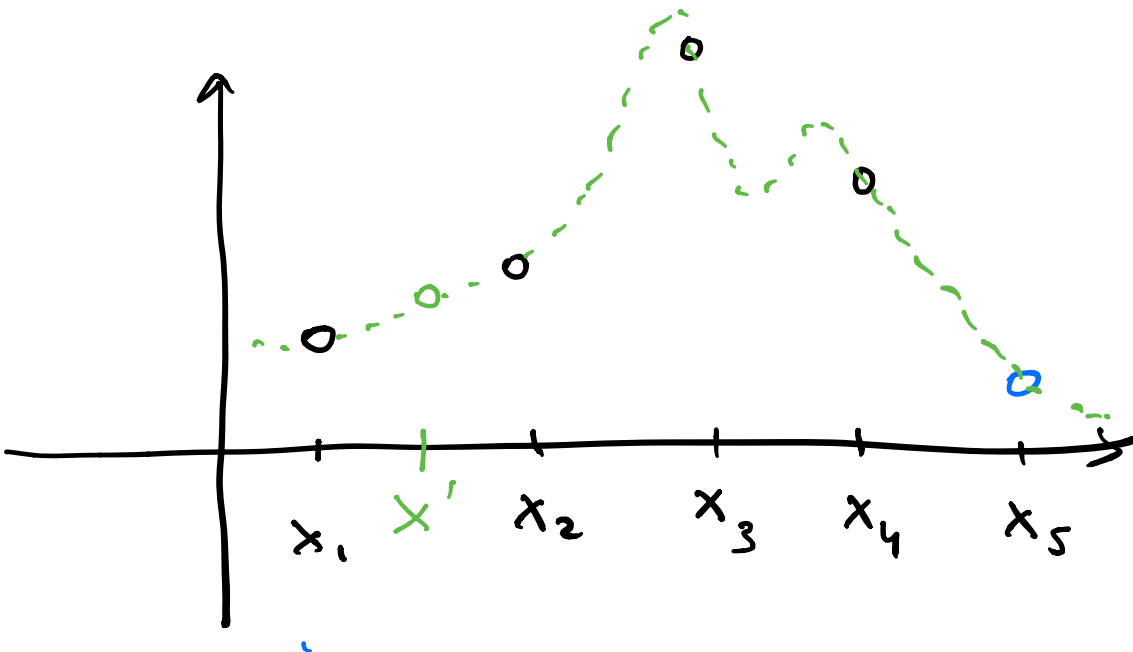
$$x_1 \rightarrow y_1$$

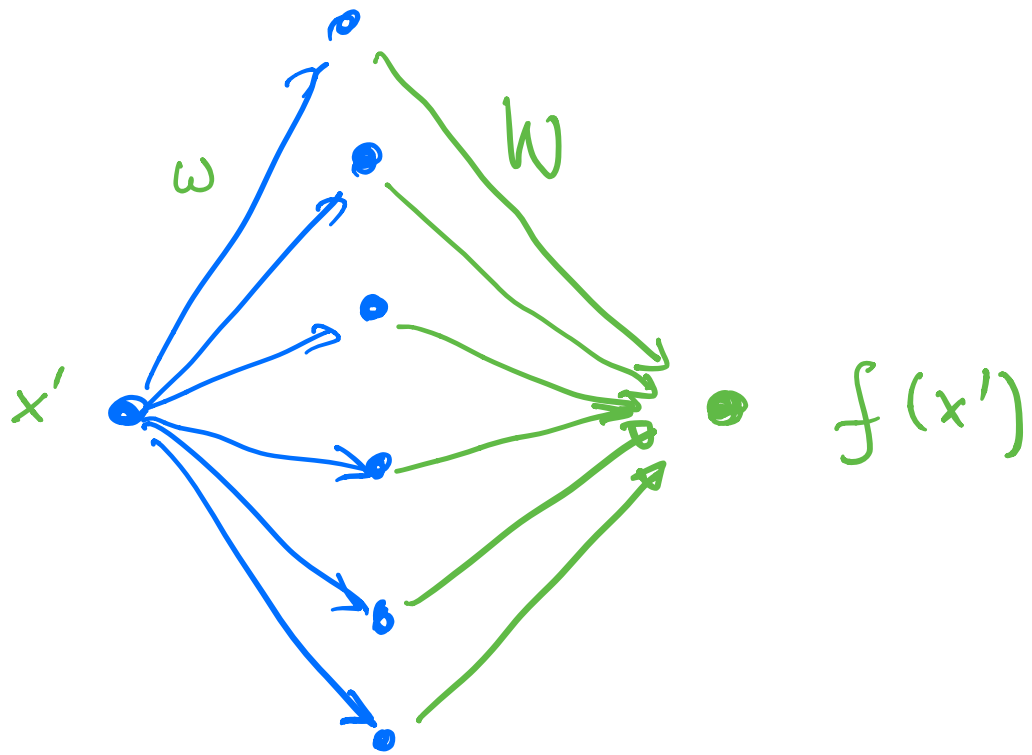
$$x_2 \rightarrow y_2$$

$$x_3 \rightarrow y_3$$

$$x_4 \rightarrow y_4$$

$$x_5 \rightarrow y_5$$





AUTO ENCODER

