

Redes neuronales frente a otras "técnicas"

Las redes neuronales tienen gran potencial en áreas como el reconocimiento de voz e imágenes, en las cuales se persiguen muchas hipótesis en paralelo. A diferencia de la arquitectura Von Neumann (proceso secuencial de instrucciones), las NN exploran varias hipótesis a la vez aprovechando el paralelismo inherente de la red.

Las NN son + robustas (tolerantes a fallos). Las NN se adaptan con el tiempo ("aprenden"). Esta cond. es clave para afrontar problemas en los q. aparecen nuevos datos constantes, muchas veces con variaciones no significativas. En cambio, los métodos estadísticos no son adaptativos: procesan todos los datos de una vez.

Las NN son más robustas cuando los datos provienen de procesos no lineales y no Gaussianos.

Las NN tienen múltiples aplicaciones:

- determinar la clase de un patrón corrupto por ruido
- mem. asociativa
- cuantización/clustering
- optimización

se aprovechan el paralelismo inherente de la red.

Las NN son + robustas (tolerantes a fallos). Las NN se adaptan con el tiempo ("aprenden"). Esta cond. es clave para afrontar problemas en los q. aparecen nuevos datos constantes, muchas veces con variaciones no significativas. En cambio, los métodos estadísticos no son adaptativos: procesan todos los datos de una vez.

Las redes neuronales ^{artificiales} se inspiran en las biológicas, pero eso no debe inducirnos a pensar q. las NN biológicas funcionan como las artificiales. He aquí algunas diferencias:

- las NN biológicas no aplican los ppos de los circuitos digitales, es decir, son analógicas. Las señales biológicas no son asincrónicas (duración variable) ni sincrónicas (reloj, control), ni tienen un umbral preciso y etc.

- Ni las neuronas ni las sinapsis son elementos de mem. biestables: no cambian de un estado a otro.

- no hay "instrucciones rígidas" ni "códigos de control"

- los circuitos cerebrales no implementan la recursividad

- las NN^{bio} son asincrónicas (en el sentido d. q. la red se reconfigura constantemente), mientras ^{en} las artificiales los pesos se actualizan de forma periódica.

- las bio aprenden con pocos ej., mientras las artificiales convergen lentas

comparación

- Ni las neuronas ni las sinapsis son elementos de mem. biestables: no cambian de un estado a otro.

- no hay "instrucciones rígidas" ni "códigos de control"

- los circuitos cerebrales no implementan la recursividad

- las NN^{bio} son asincrónicas (en el sentido d. q. la red se reconfigura constantemente), mientras ^{en} las artificiales los pesos se actualizan de forma periódica.

- las bio aprenden con pocos ej., mientras las artificiales convergen lentas

REDES NEURONALES

Una red neuronal artificial es un sistema de computación desarrollado como una generaliz. de los modelos matemáticos del conocimiento humano (neurobiología), por eso presenta muchas similitudes con las redes neuronales biológicas:

- red formada por muchas unids. de proceso simples, q. tratan la info. localiz (paralelismo masivo)
- cada neurona recibe muchas señales de entrada
- cada señal puede ser modificada por un peso en la sinapsis
- la neurona combina las señales de entrada ponderadas para producir una señal de salida, q. a su vez puede propagarse a una o varias neuronas
- la mem. es distribuida: a largo plazo reside en los pesos, a corto plazo en las señales
- los pesos pueden alterarse por la experiencia
- los neurotransmisores pueden ser excitadores o inhibidores

des con las redes neuronales biológicas.

- red formada por muchas unids. de proceso simples, q. tratan la info. localiz (paralelismo masivo)
- cada neurona recibe muchas señales de entrada
- cada señal puede ser modificada por un peso en la sinapsis
- la neurona combina las señales de entrada ponderadas para producir una señal

Las redes neuronales artificiales (NN) se caracterizan por

- arquitectura $\rightarrow$ n° de capas $\rightarrow$ monocapa $\rightarrow$ multicapa

interconexiones

flujo directo (feed forward)

sólo conexiones de la capa N a la $N+1$

~~recurrentes~~ / ~~iterativas~~ (feedback) ~~con~~ retroalimentación

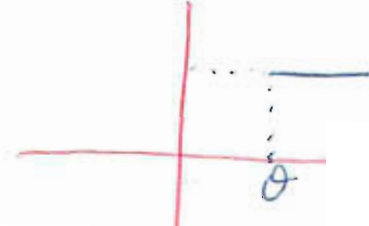
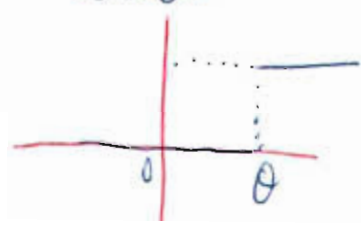
$\rightarrow$ conexiones entre N y $N+1$ y entre $N+1$ y N

- algoritmos de aprendizaje $\rightarrow$ supervisado $\rightarrow$ no supervisado

- algoritmos $\rightarrow$ iterativos / recurrentes: van refinando la solución
sementales: cada paso se aplica 1 sola vez

- func. de activación: sirven para restringir el rango de valores de una señal y para introducir la "no linealidad" en la red (ya q. una combinación de func. lineales es lineal, y puede ser insuficiente para q. el modelo sea fiel)

- identidad: $f(y_j) = y_{enj} = y_j$
- escalón



$\theta = \text{umbral}$

~~recurrentes~~ / ~~iterativas~~ (feedback) ~~con~~ retroalimentación

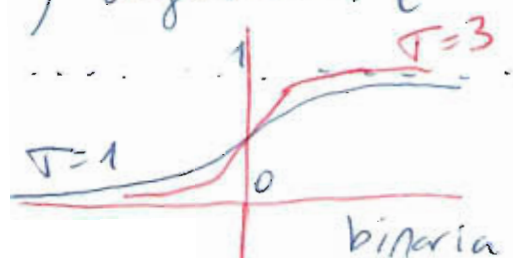
$\rightarrow$ conexiones entre N y $N+1$ y entre $N+1$ y N

- algoritmos de aprendizaje $\rightarrow$ supervisado $\rightarrow$ no supervisado

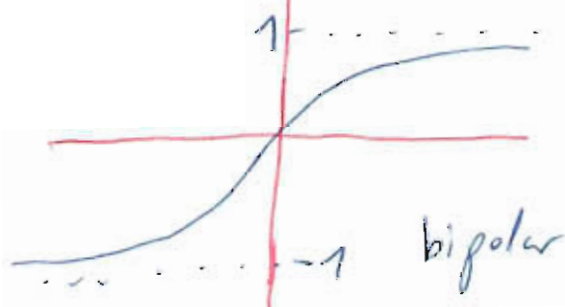
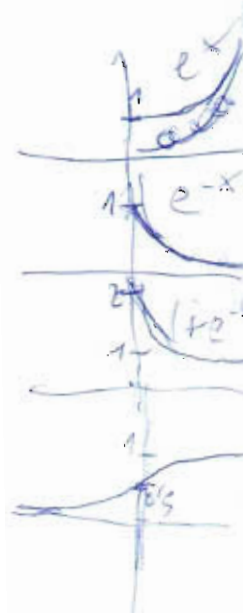
- algoritmos $\rightarrow$ iterativos / recurrentes: van refinando la solución
sementales: cada paso se aplica 1 sola vez

- func. de activación: sirven para restringir el rango de valores de una señal y para introducir la "no linealidad" en la red (ya

- sigmoidal: suele utilizarse pq. es acotada y diferenciable



$$f(x) = \frac{1}{1 + e^{-x}}$$



$$g(x) = 2f(x) - 1 = \frac{1 - e^{-x}}{1 + e^{-x}}$$

Algunas arquitecturas emplean la derivada de la func. de activación:

$$[f'(x)] \Rightarrow \frac{d}{dx} u^n = n u^{n-1} \frac{du}{dx} \quad u = 1 + e^{-x} \quad \frac{du}{dx} = -e^{-x} \quad n = -1$$

$$\rightarrow -1 \cdot (1 + e^{-x})^{-2} \cdot (-e^{-x}) = \frac{e^{-x}}{(1 + e^{-x})^2}$$

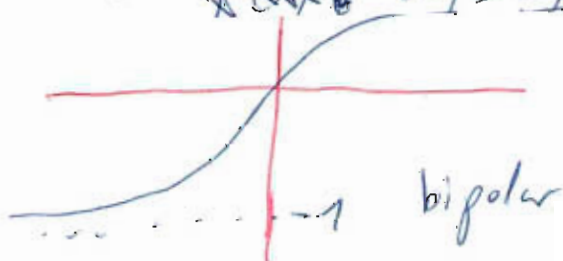
La derivada se puede expresar a partir de la propia $f(x)$:

$$f'(x) = \frac{e^{-x}}{(1 + e^{-x})^2} = \frac{1}{(1 + e^{-x})} \cdot \frac{e^{-x}}{(1 + e^{-x})} = f(x) \cdot (1 - f(x))$$

Normalmente, $\sigma = 1$, por lo q.

$$f'(x) = f(x) [1 - f(x)]$$

$$[g'(x)] \Rightarrow \frac{d}{dx} u^n = n u^{n-1} \frac{du}{dx}$$



$$g(x) = \frac{1 - e^{-x}}{1 + e^{-x}} \quad g'(x) = \frac{e^{-x}}{(1 + e^{-x})^2} = \frac{1 - e^{-x}}{1 + e^{-x}}$$

Algunas arquitecturas emplean la derivada de la func. de activación:

$$[f'(x)] \Rightarrow \frac{d}{dx} u^n = n u^{n-1} \frac{du}{dx} \quad u = 1 + e^{-x} \quad \frac{du}{dx} = -e^{-x} \quad n = -1$$

Clasificación redes

- McCulloch - Pitts
- Aprendizaje no supervisado
- Feedback¹ (flujo total)

- ART
- Hopfield
- BAM
- SOM

Aprendizaje competición: Hamming, Maxnet, sombrero mexicano,

- Feedforward-only² (flujo directo)
- Contrapropagación (CP)

- Aprendizaje supervisado

- Feedback¹
- BSB, hetero/asociativistas

- Feedforward-only²
- Perceptron
- Adaline / Madaline
- Retropropagación (BP)
- LVQ
- Hebb

- SOM

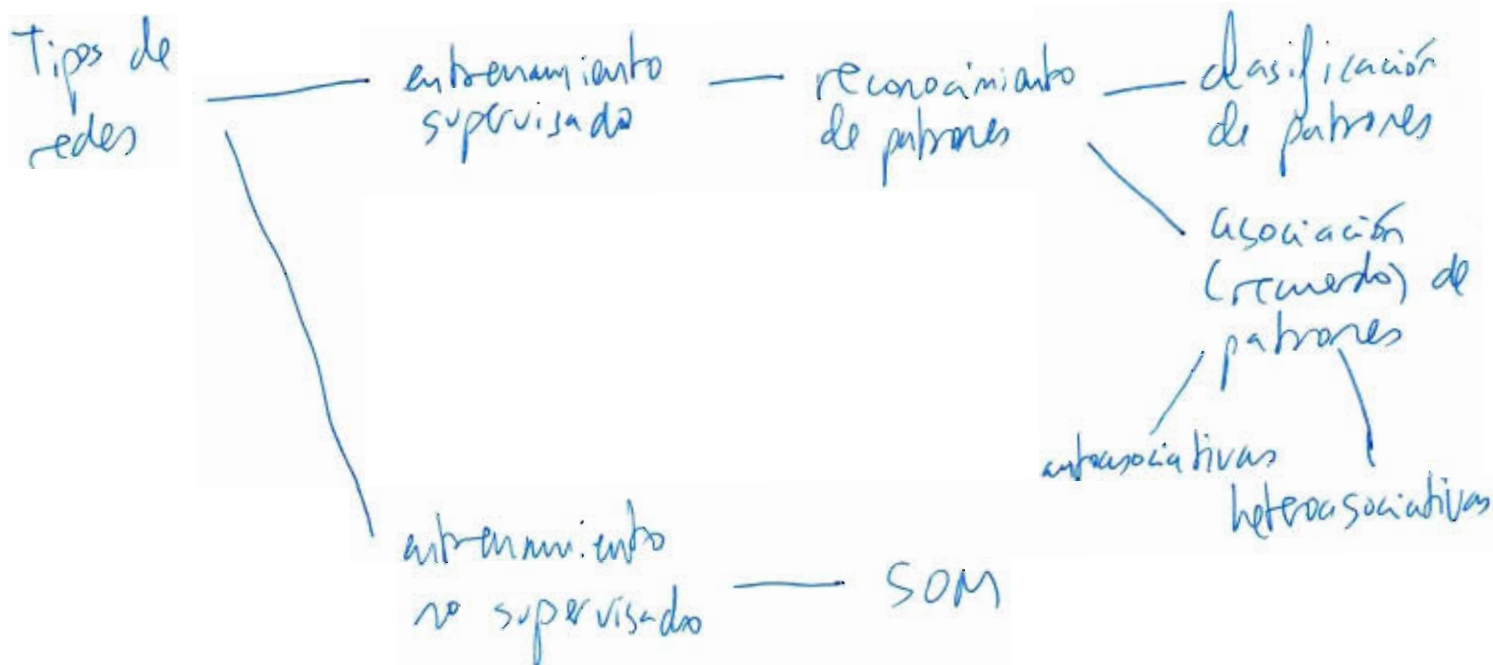
Aprendizaje competición: Hamming, Maxnet, sombrero mexicano,

- Feedforward-only² (flujo directo)
- Contrapropagación (CP)

- Aprendizaje supervisado

- Feedback¹
- BSB, hetero/asociativistas

- Feedforward-only²



Conjunto entrenamiento → ajuste pesos
→ comprobar aprendizaje

Conjunto de prueba → verificar aprendizaje independiente del de entrenamiento

Ambos deben ser significativos (no suficiente ej.) y representativos (diversos)

- reconocimiento** **Nueva entrada → clases**
- Clasif. de patrones: cada patrón (vector de entrada) pertenece (o no) a una clase (categoría). Se conocen las clases de los patrones del conjunto de entrenamiento. Puede haber varias clases (lo q. conlleva varias unids. de salida → unids. salida = n° clases salida)
 - En estas redes el bias y el umbral son equivalentes, por lo no aporta nada incluir ambos. Sin embargo, no
 - entrenamiento no supervisado — SOM

Conjunto entrenamiento → ajuste pesos
→ comprobar aprendizaje

Conjunto de prueba → verificar aprendizaje independiente del de entrenamiento

Ambos deben ser significativos (no suficiente ej.) y representativos

lineal de fues. es lineal)

→ Mapco patrón entrada → patrón salida

- Asoc. de patrones: la red almacena un conjunto de asociaciones entre patrones de entrada (s) y salida (t).

Si $s=t$, la red es autoasociativa. Si $s \neq t$, heteroasociativa.

Cuando se usa la regla de Hebb, suelen normalizarse los pesos finales con un factor $1/n$ ($n = n^\circ$ unids. de la red)

- Competición: sólo una de las neuronas gana

- Autoorganizativas: la unid. cuyo peso esté más próximo al vector de entrada es la q. se actualiza. Para calcular la proximidad se puede usar el cuadrado de la distancia euclídea o el producto vectorial (correlación)

Solución problemas no linealmente separables

- Cambio codif. patrones (p.e., añadir una ~~o~~ neurona oculta q. ayude a separar los patrones)

Cond. de parada

- n° tipo de épocas

- el error disminuye por debajo de un umbral

- la variación en los pesos es irrelevante

... $1/n$ n° de unids. de la red

- Competición: sólo una de las neuronas gana

- Autoorganizativas: la unid. cuyo peso esté más próximo al vector de entrada es la q. se actualiza. Para calcular la proximidad se puede usar el cuadrado de la distancia euclídea o el producto vectorial (correlación)

REGLAS DE APRENDIZAJE

Regla de Hebb (Donald Hebb, 1949)

"Cuando un axón de una célula A está lo suficiente cerca para excitar una célula B, y toma parte repetida en el proceso de disparo de dicha célula, se produce algún tipo de cambio metabólico, en una de las células (o en las 2) q. hace q. la eficacia con la q. A dispara a B se vea incrementada"

Es decir, cuando 2 neuronas se activan repetidas al mismo tiempo, tienden a asociarse, de modo q. la actividad en una incita la actividad de la otra. En pocas palabras, se refuerzan las conexiones de las neuronas q. se activan simultáneamente.

Hebb propuso como ej. de aumento en la sinapsis el aumento del n° de terminales sinápticos o su alargamiento.

La regla de Hebb como tal es bastante genérica y ha dado pie a numerosas formulaciones. La más simple es:

$$[1] \quad w_{ij} = x_i \cdot x_j \xrightarrow{\text{binaria}} \begin{matrix} 0 \cdot 0 = 0 \\ 0 \cdot 1 = 0 \\ 1 \cdot 0 = 0 \\ 1 \cdot 1 = 1 \end{matrix}$$

duce algún tipo de cambio metabólico, en una de las células (o en las 2) q. hace q. la eficacia con la q. A dispara a B se vea incrementada"

Es decir, cuando 2 neuronas se activan repetidas al mismo tiempo, tienden a asociarse, de modo q. la actividad en una incita la actividad de la otra. En pocas palabras, se refuerzan las conexiones de las neuronas q. se activan simultáneamente.

Eso, cuando el refuerzo también se produce cuando varias neuronas se desactivan a la vez, se habla de regla de Hebb extendida.

Una formulación + biológica sería:

$$w(t+1) = w(t) + (a_{pe} \cdot a_{ps}) \quad [2]$$

(activación presináptica/postsináptica)

la formulación + habitual hoy en día es

$$\Delta w_{ij} = x_i y_j \quad [3]$$

~~q. también se puede plantear~~

En las redes asociativas se aplica la fórmula [1] vectorialmente:

$$w = \bar{x}^T \cdot \bar{y} \quad [4]$$

En el perceptrón se aplica una regla equivalente (cuando la salida obtenida no coincide con la deseada):

$$\Delta w_i = \alpha t x_i$$
$$w(t+1) = w(t) + (a_{pe} \cdot a_{ps}) \quad \dots$$

(activación presináptica/postsináptica)

la formulación + habitual hoy en día es

$$\Delta w_{ij} = x_i y_j \quad [3]$$

~~q. también se puede plantear~~

Regla Delta: busca minimizar la diferencia entre el valor real producido en la capa de salida de la red y el valor esperado u objetivo:

$$\Delta w_{ij} = \alpha (t_j - y_{inj}) x_i$$

$t \rightarrow$ valor esperado

$y_{in} \rightarrow$ valor obtenido (antes de aplicar la func. de activación)

$x_i \rightarrow$ valor de entrada

$\alpha \rightarrow$ tasa de aprendizaje

- grande $\rightarrow$ convergencia rápida, con riesgo de perder mínimos u oscilar entre ellos
- peg. $\rightarrow$ convergencia lenta y segura

La regla delta minimiza la func. de error. ~~por~~ Para acelerar el proceso se puede recurrir a técnicas de optimización no lineal. En el caso del perceptrón, suele utilizarse el método de descenso del gradiente.

$t \rightarrow$ valor esperado

$y_{in} \rightarrow$ valor obtenido (antes de aplicar la func. de activación)

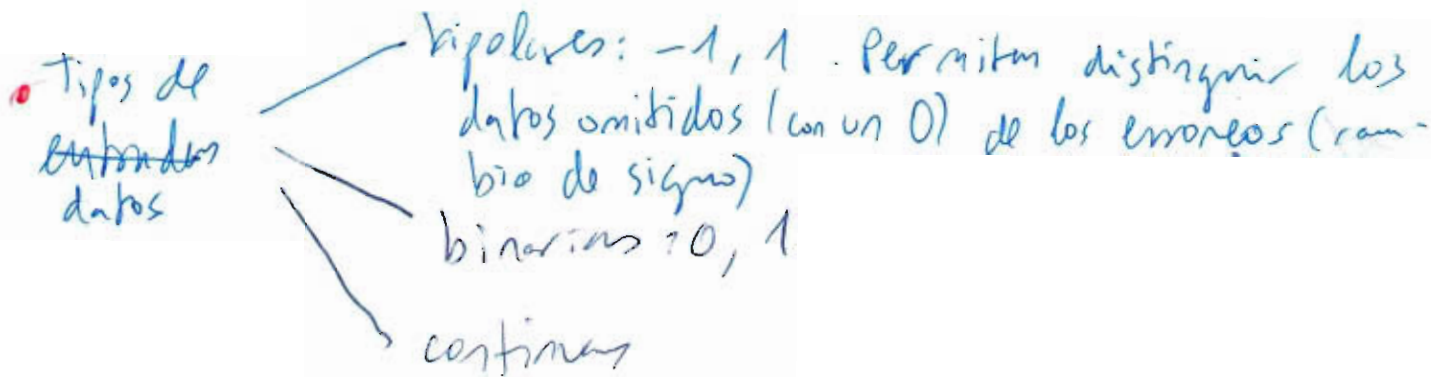
$x_i \rightarrow$ valor de entrada

$\alpha \rightarrow$ tasa de aprendizaje

- grande $\rightarrow$ convergencia rápida, con riesgo de perder mínimos u oscilar entre ellos

Aprendizaje Kohonen

$$\Delta w_{ij}(t) = \alpha (x_i - w_{ij}(t-1))$$



Aprendizaje Kohonen $\Delta w_{ij} = \alpha (x - w_{jold})$

Regla aprendizaje perceptrón (equivale a Hebb, salvo q. se aplica sólo cuando hay fallos)

$$w_{i(new)} = w_{i(old)} + \alpha t x_i$$

t = valor esperado

x_i = entrada

α = tasa aprendizaje

→ regla delta generalizada

FUNCS. DE ACTIVACIÓN

Hebb

- red de Hebb (B)
- heteroasociativa (E)
- autoasociativa (E)
- BSB (I-E)
- Hopfield (E)
- BAM (E)

Pelta

- Adaline (I,E) (B)
- Madaline (E) (B)
- (heteroasociativa)
- (autoasociativa)
- (contemporáneas)
- (retropropagación) (S)

Perceptron — Perceptron (E)(B)

Otros

- Maxnet (F)
- Sombrero mej. como (E-E)
- Hanning (B)
- ART
- McClelland-Pitts (E)

Kohonen — LUG
SOM

Regla aprendizaje perceptrón (equivale a Hebb, salvo q. se aplica sólo cuando hay fallos)

$$w_{i(new)} = w_{i(old)} + \alpha t x_i$$

t = valor esperado

x_i = entrada

α = tasa aprendizaje

→ regla delta generalizada

FUNCS. DE ACTIVACIÓN

Hebb

- red de Hebb (B)
- heteroasociativa (E)

Pelta

- Adaline (I,E) (B)
- Madaline (E) (B)
- (heteroasociativa)

REGLAS DE APRENDIZAJE

Limitaciones: codif. bipolar E/S { separ. actividad inactiva }
 Neuronas limitadas

$$\Delta w_i = x_i \cdot y$$

Regla de Hebb (1949)

→ mecanismo local y dependiente del tiempo

se refuerzan las conexiones de las neuronas q. se activan simultáneamente. No se usa en la actualidad. "Cuando un axón de la célula A está lo bastante próximo a la célula B como para excitarla, y de forma repetida o persistente colabora en su activación, tienen lugar algunos procesos de colapso o cambio metabólico en una o ambas células, tal q. la eficacia de A, con respecto a B, se incrementa."

Regla de Hebb extendida: también se refuerzan las conexiones cuando varias neuronas se desactivan a la vez.

dependiendo de la codif. se puede llegar a unos pesos q. separan linealmente (o no) los datos, por eso a veces se cambia de codif.

$$w_{ij}(\text{new}) = w_{ij}(\text{old}) + x_i \cdot y_j \quad (\bar{x} \cdot \bar{y} \equiv \text{correlación})$$

Si los vectores x_i no son ortogonales, la respuesta de la neurona incluye interferencias

Regla Delta: minimiza la diferencia entre la entrada (o "mínimos cuadrados")

red de la unid. de salida (y_{in}) y el valor objetivo (t)

$$w_i(\text{new}) = w_i(\text{old}) + \alpha (t - y_{\text{in}}) x_i$$

α grande → convergencia rápida con riesgo de perder mínimos
 α pequeña → convergencia lenta y segura

La regla de Hebb se puede calcular en varios pasos (tanto como palabras en almacenamiento). Ej.

$$P1: (-1, 1, -1, 1, -1, 1) \quad t = (-1, 1)$$

$$P2: (-1, 1, 1, 1, -1, -1) \quad t = (1, 1)$$

$$W_1 = s_1^T \cdot t_1 = \begin{pmatrix} 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \end{pmatrix}$$

$$W_2 = s_2^T \cdot t_2 = \begin{pmatrix} -1 & -1 \\ 1 & 1 \\ 1 & 1 \\ 1 & 1 \\ -1 & -1 \\ -1 & -1 \end{pmatrix}$$

$$W = \begin{pmatrix} 0 & -2 \\ 0 & 2 \\ 2 & 0 \\ 0 & 2 \\ 0 & -2 \\ -2 & 0 \end{pmatrix} \quad W_1 + W_2 = \begin{pmatrix} 0 & -2 \\ 0 & 2 \\ 2 & 0 \\ 0 & 2 \\ 0 & -2 \\ -2 & 0 \end{pmatrix}$$

esta abierta con varias formulaciones. La + simple es $\Delta w_{ij} = x_i \cdot y_j$

cuando varias neuronas se desactivan a la vez, dependiendo de la codif. se puede llegar a unos pesos q. separan linealmente (o no) los datos, por eso a veces se cambia de codif.

$$w_{ij}(\text{new}) = w_{ij}(\text{old}) + x_i \cdot y_j \quad (\bar{x} \cdot \bar{y} \equiv \text{correlación})$$

Si los vectores x_i no son ortogonales, la respuesta de la neurona incluye interferencias

Regla Delta: minimiza la diferencia entre la entrada (o "mínimos cuadrados")

red de la unid. de salida (y_{in}) y el valor objetivo (t)

$$w_i(\text{new}) = w_i(\text{old}) + \alpha (t - y_{\text{in}}) x_i$$

α grande → convergencia rápida con riesgo de perder mínimos
 α pequeña → convergencia lenta y segura

$$y = \sum x_i \cdot w_i$$

evol. redes neuronales

Redes básicas

→ McCulloch

Implementar
funciones básicas
con NN
a partir de
valores de
ej. capacitancia
resistencias
fijos. Vm

una implementación
de la func. se puede
aplicar a nuevas
redes

Hebb

Perceptrón

Adaline

Madaline

~~pesos fijos~~

pesos fijos
umbral de cada neurona

→ aprendizaje (Hebb)
(clasif. patrones)

tasa aprendizaje
func. activación (umbral)

multicapa (opcional)
se pasan varias veces los
patrones (iteraciones)

varias salidas
aprendizaje delta

→ Capa oculta

backpropagation
capa oculta
propagando error

asoc. patrones

Asociativas iterativas

→ Heteroasociativa → inicializ. pesos Hebb

Se parte de un patrón
incompleto y se quiere
saber a q. patrón memo-
rizado se parece +

↓

Autoasociativa

→ $x = y$

→ vectores ortogonales

Algoritmo con func.
umbral

→ repetición

$$x_{in}(i+1) = x(i)w$$

→ θ_i

hijos. Vm
una implementación
de la func. se puede
aplicar a nuevas
redes

Perceptrón

Adaline

Madaline

tasa aprendizaje
func. activación (umbral)

multicapa (opcional)
se pasan varias veces los
patrones (iteraciones)

varias salidas
aprendizaje delta

→ Capa oculta

backpropagation
capa oculta
propagando error

asoc. patrones

Competitivos

Maxnet
elegir unidades + activadas
Sombrero mexicano

→ neurona + activada
pesos fijos
no hay aprendizaje
→ radio, pesos refuerzo e inhibición

Hamming

→ pesos fijos elige la unidad.
mayor número se potencia + al de ~~entrada~~ entrada

asociación / clasificación / reconocimiento
caracteres, palabras, imágenes

SOM

los pesos representan al vibrante ^{radio}
topología y vecindad
aprendizaje Kohonen
no supervisado
elige neurona + similar y
sus vecinas por refuerzo
el n.º de clusters se estima
tasa aprendizaje y radio de vecindad
sin topología / vecindad
se conocen las clases (o clusters)
si la clase coincide, refuerza;
si no, inhibe

LVQ

Contrapropagación

sin topología / pesos fijos elige la unidad.
mayor número se potencia + al de ~~entrada~~ entrada

SOM

los pesos representan al vibrante ^{radio}
topología y vecindad
aprendizaje Kohonen
no supervisado
elige neurona + similar y

Redes Neuronales Artificiales – resumen características

Red	Aprendizaje	Tipo		Codif		Func activ	Bias	Entren.	Iterat.	Inic.	Otros
McCulloch-Pitts	No			Binaria		1 $y \geq \theta$ 0 $y < \theta$	No	No	No	No	θ_i
Hebb	Hebb	Clasif. supervisada		Bipolar		No	Sí	Sí	Sí	$W(0)=0$	
Perceptrón						1 $y > \theta$ 0 $-\theta \leq y \leq \theta$ -1 $y < -\theta$					
Adaline	Delta					$Y=x$ (entren.) 1 $y \geq 0$ -1 $y < 0$ (apl.)				$W(0)=aleat.$	
Madaline						1 $y \geq 0$ -1 $y < 0$					
Retropropagación	Delta generaliz.			Binaria / Bipolar		Sigmoidal					
Heteroasociativa	Hebb	Asociativas	Heteroasoc.	Bipolar	Convers.	1 $y > 0$ 0 signif. $Y(t) y_{in}(t-1) = \theta$	No	No	No	$W(0)=\sum s't$	
BAM						-1/0 $y \leq 0$			Sí		θ_j
Autoasociativa			Autoasoc.		-	0 signif.			No	$W(0)=\sum s's$	
Autoasoc. Con func. Umbral									Sí	$W_{ij} = 0$	
Hopfield				Convers.							
Maxnet	No	Competitivas		Continua		$X x > 0$ 0 $x \leq 0$ $X_{max} x > x_{max}$ $X 0 \leq x \leq x_{max}$ 0 $x < 0$				$W(0)=-\epsilon / 1$	
Sombrero mejicano											
Hamming											
SOM	Kohonen			Bipolar		No	$B(0)=n/2$	Sí	No	$W(0)=e/2$	>act.
LVQ				Cualquiera			No		Sí	$W(0)=aleat$	<dist.
Contrapropagación										$W(0)=aleat$	
ART1	"ART"			Binaria						$W(0)=aleat$	
										$B(0) < L/(L-1+n)$	>act.
Adaline	Delta					-1 $y < \theta$ $Y=x$ (entren.) 1 $y \geq 0$ -1 $y < 0$ (apl.)				$W(0)=aleat.$	
Madaline						1 $y \geq 0$ -1 $y < 0$					
Retropropagación	Delta generaliz.			Binaria / Bipolar		Sigmoidal					
Heteroasociativa	Hebb	Asociativas	Heteroasoc.	Bipolar	Convers.	1 $y > 0$ 0 signif. $Y(t) y_{in}(t-1) = \theta$	No	No	No	$W(0)=\sum s't$	
BAM						-1/0 $y \leq 0$			Sí		θ_j
Autoasociativa			Autoasoc.		-				No	$W(0)=\sum s's$	

Entrenamiento / Iteraciones

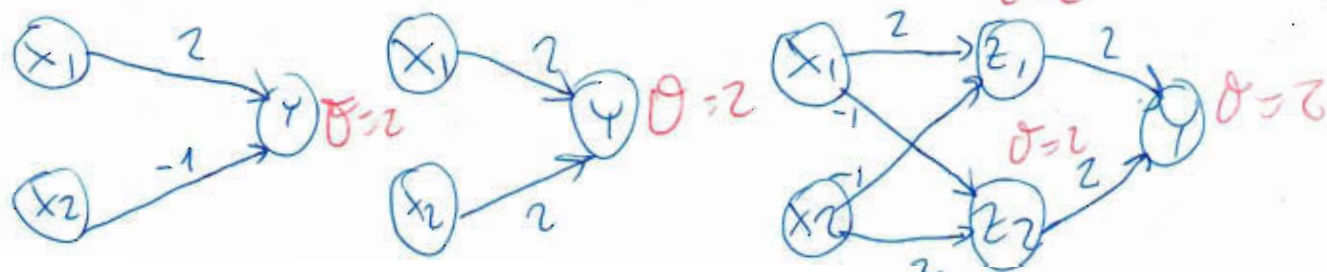
Red	Codif.	Func. activ.	Bias	Tipo	Obs
McCulloch-Pitts	Binaria	$1 \ y \geq 0$ $0 \ y < 0$	Si No		θ_i NE, NI
Hebb	Bipolar	—	SI	Clasif. sep.	Hebb E, I, F E, I $w(0) = 0$
Perceptron	(salida)	$1 \ y > 0$ $0 \ -0 \leq y \leq 0$ $-1 \ y < 0$	tty		
Adaline		$y = x$ (brain) $1 \ y \geq 0$ (app.) $-1 \ y < 0$ (app.)			Delta E, F $w(0) = \text{aleat.}$
Madaline		$1 \ y \geq 0$ $-1 \ y < 0$	tty		E, I, F , w/o aleat (generalizado)
Backprop.	Binaria Bipolar	Sigmoidal	$0'10 \leq P(N+M)+$ $+M \leq 0'150$		
Heterosoc.		$1 \ y > 0$ $-1/0 \ y \leq 0$	No	Asociat. (memorias)	θ_i Hebb $(w_{ij} = \sum s_i \cdot t_j)$ $(w_{ii} = 0)$ $w(0) = \sum s_i \cdot s_i$
Auto-soc.	Bipolar	Madaline		NE	
Auto. con func. umbral		0 importante		I	
Hopfield	(conversion)			I	
RAM Perceptron	(conversion) (salida)	$y(t), y_n(t-1) = 0$ $0 \ -0 \leq y \leq 0$ $-1 \ y < 0$	tty		$(w_{ij} = \sum s_i \cdot t_j)$ $w(0) = 0$
Adaline		$y = x$ (brain) $1 \ y \geq 0$ (app.) $-1 \ y < 0$ (app.)			Delta E, F $w(0) = \text{aleat.}$
Madaline		$1 \ y \geq 0$ $-1 \ y < 0$	tty		E, I, F , w/o aleat (generalizado)

McCulloch - Pitts (1943)

TIPO:

USOS:

ARQUITECTURA



AND-NOT

OR

XOR (= ANDNOT OR ANDNOT)

w = peso positivo $\rightarrow$ camino excitatorio

$-p$ = " negativo $\rightarrow$ " inhibidor

Todos los caminos excitatorios tienen los mismos pesos

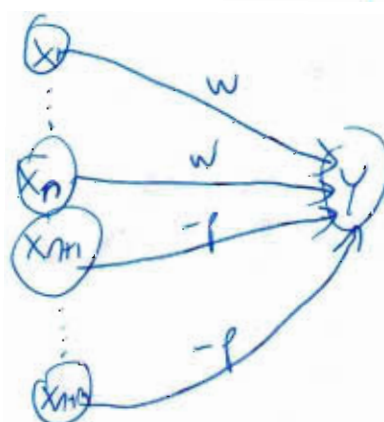
$$y(t) = z_1(t-1) \text{ OR } z_2(t-1) = x_1(t-2) \text{ AND NOT } x_2(t-2) \text{ OR } x_2(t-2) \text{ AND NOT } x_1(t-2)$$

ALGORITMO

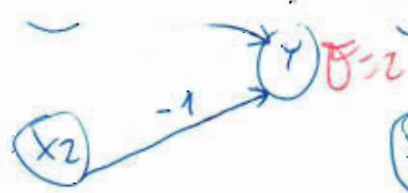
Pesos: fijos (positivos o neg.)

Extremos: no hay

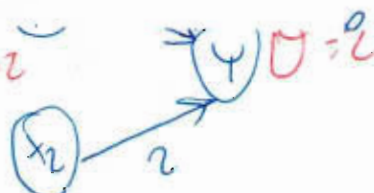
Func. activación



1: —

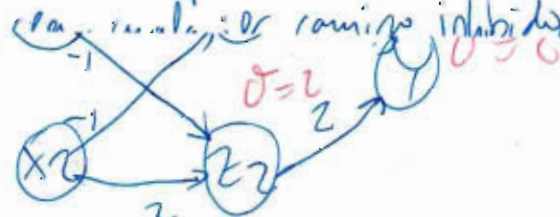


AND-NOT



OR

$\theta > nW - p$ (inhibición absoluta, camino inhibidor $\neq 0$)



XOR (= ANDNOT OR ANDNOT)

w = peso positivo $\rightarrow$ camino excitatorio

$-p$ = " negativo $\rightarrow$ " inhibidor

Todos los caminos excitatorios tienen los mismos pesos

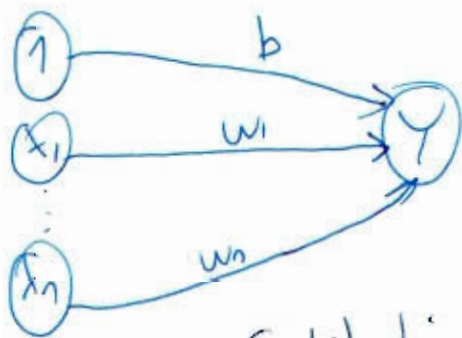
$$y(t) = z_1(t-1) \text{ OR } z_2(t-1) = x_1(t-2) \text{ AND NOT } x_2(t-2) \text{ OR } x_2(t-2) \text{ AND NOT } x_1(t-2)$$

HEBB

TIPO: Clasif. patrones (ent. supervisado)

USOS:

ARQUITECTURA



Codif. hipotesis entrada y salida

ALGORITMO:

Inic. pesos: $w_i = 0 \quad i \in \{1, n\}$ ($w_0 = b$)

• Regla aprendizaje: Hebb extendida

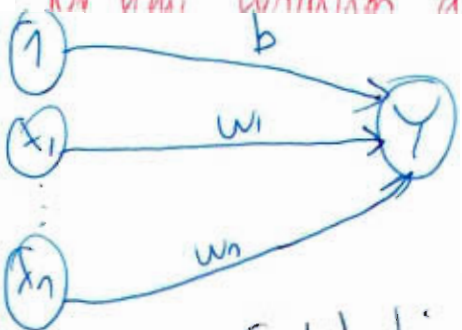
$$w_i(\text{new}) = w_i(\text{old}) + x_i \cdot t$$

$$b(\text{new}) = b(\text{old}) + t \quad (x_0 = 1)$$

$$\begin{cases} x = s \\ y = t \end{cases}$$

Ajustar pesos y bias

En una condición de record (los patrones se presentan)



Codif. hipotesis entrada y salida

PERCEPTRON

TIPO: clasif. patrones (ant. supervisado), Feed Forward (FF)

USOS: clasif. automática

ARQUITECTURA: = Hebb

x_i : entradas

t : valor esperado de y

y : respuesta unid. salida

ALGORITMO ENTRENAMIENTO

Inic. pesos: $w_i = 0, i \in [0, N]$

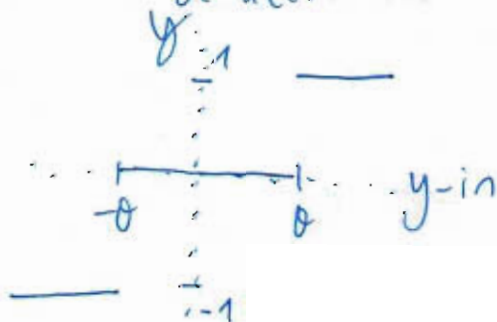
Tasa aprendizaje $0 < \alpha \leq 1$

Regla aprendizaje: (Hebb)

$$w_i(\text{new}) = w_i(\text{old}) + \alpha t x_i$$

Entrada de y : $y_{\text{in}} = \sum_{i=0}^n x_i w_i$

Func. de activación



$$y = \begin{cases} 1 & y_{\text{in}} > 0 \\ 0 & -\theta \leq y_{\text{in}} \leq \theta \\ -1 & y_{\text{in}} < -\theta \end{cases}$$

Cond. de parada: los pesos no han cambiado en k

Último error

t : valor esperado de y

y : respuesta unid. salida

ALGORITMO ENTRENAMIENTO

Inic. pesos: $w_i = 0, i \in [0, N]$

Tasa aprendizaje $0 < \alpha \leq 1$

Regla aprendizaje: (Hebb)

$$w_i(\text{new}) = w_i(\text{old}) + \alpha t x_i$$



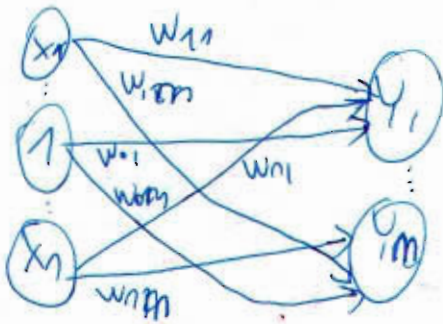
APALINE

ADaptive Linear ~~Neuron~~ NEURON
(Widrow & Hoff, 1960)

TIPO: clasif. patronal. supervisado, FF

USOS:

ARQUITECTURA



$$y_{inj} = \sum_{i=0}^n x_i w_{ij}$$

Q patrones

Codif. bipolar ←

ALGORITMO ENTRENAMIENTO

Inic. pesos: valores aleatorios peg.

α → muy grande: el aprendizaje no converge

→ muy peg.: el " " es muy lento

para n=1, 0.1 ≤ α ≤ 1

n = n° entrenamientos

Regla aprendizaje:

Delta ←

Se aplica a todos los entrenamientos

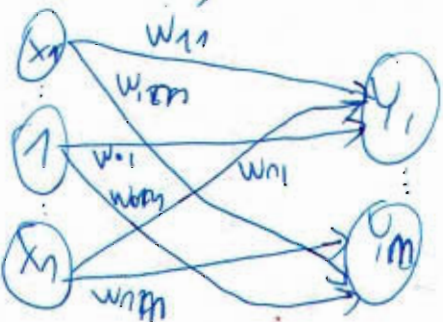
$$\Delta w_{ij} = \alpha (t_j - y_{inj}) x_i$$

$$w_{ij}(new) = w_{ij}(old) + \Delta w_{ij}$$

$$y_{in} = b + \sum x_i w_i$$

Func. activaci3n:

y



$$y_{inj} = \sum_{i=0}^n x_i w_{ij}$$

Q patrones

Codif. bipolar ←

$$y_{in} = \bar{x} \cdot \bar{w}$$

$$y = f(y_{in})$$

ALGORITMO ENTRENAMIENTO

Inic. pesos: valores aleatorios peg.



Mary Adaline

$$Y = z_1 \text{ OR } z_2 \quad (v_1 = v_2 = b_3 = \frac{1}{12})$$

(Hb. se pueden utilizar otras func. como ANP = "mayoría")

ANP o "majoria")
 "unids. de un capr correctam con
 "de la "sgte)

MRF (solo se ajustan los pesos w_{ij})

ω_{ij} = valor aleatorio peg.

$\Delta \geq$ valor p.g.

$$\Gamma x_i = \xi_i$$

$$Z_{in} = b + \sum_j w_j$$

$$z_i = f(z - in_i)$$

$$y_1 = b_3 + z_1 v_1 + z_2 v_2$$

$$y = f(y_{-i})$$

Si $y = t \rightarrow$ salida correcta. No se modifican pesos
 loc si $t = 1$, actualizar pesos w_{ij} en la unid. j con
 entrada $z_{ij} + \text{proxima a } 0$ (2)

Si $t = -1$, actualizar pesos en unids. k con entrada $z_{in,k} > 0$

Cond. parada: variación de los pesos es aceptable,
Codif. bipolar (todas las unid. de la cap. concuerdan con la sigte)

Codex bipolar

MRF (solo se ajustan los pesos w_{ij})

u_{ij} = valor aleatório peg.

$\Delta \geq$ valor p/q.

$$\Gamma x_i = s_i$$

MRII (se ajustan todos los pesos)
Inicializar pesos
" 2

Calcular γ como en algoritmo MRI
Determinar error y actualizar pesos:
[Si $y \neq t$,
 $V_j / Z_j \in (-0.25, 0.25)$
 $Z'_j = -1 \cdot Z_j$
 Recalcular γ
 Si se ha reducido el error, actualizar
 los W_{ij} (aplicar la regla delta con Z'_j)
Condición parada: igual q. MRI

$V_j / Z_j \in (-0.25, 0.25)$
[$Z'_j = -1 \cdot Z_j$
 Recalcular γ
 Si se ha reducido el error, actualizar
 los W_{ij} (aplicar la regla delta con Z'_j)
Condición parada: igual q. MRI

RETROPROPAGACIÓN (Backpropagation)

TIPO: aprendizaje supervisado, feed forward

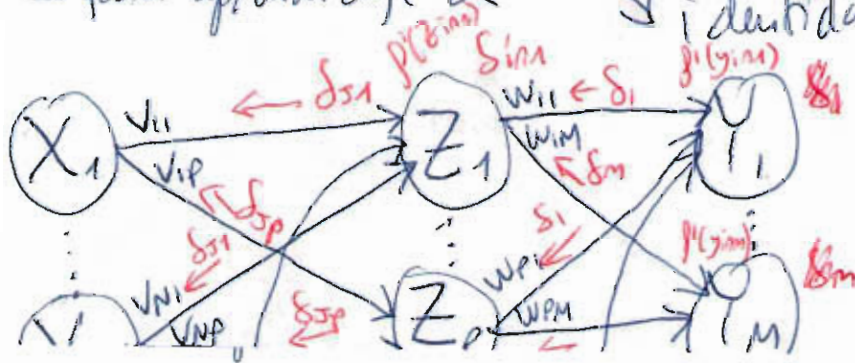
USOS: $\rightarrow$ aproximador universal (en particular, relaciones no lineales)
 $\rightarrow$ filtrado de ruido
 $\rightarrow$ compresión

Codif. binaria o bipolar

ARQUITECTURA

Perceptrón multicapa. Para cada problema se diseña una red adaptando los sigtes. parám.

- n° entradas (N)
- n° salidas (M)
- n° neuronas capa oculta (P_c)
- n° de capas ocultas (C). Una capa oculta es suficiente en la mayoría de los problemas, por lo q. trabajaremos con $C=1$ para simplificar
- func. de activación $\rightarrow$ sigmoide $\rightarrow$ binaria
- tasa aprendizaje α $\rightarrow$ identidad $\rightarrow$ bipolar



ARQUITECTURA

Perceptrón multicapa. Para cada problema se diseña una red adaptando los sigtes. parám.

- n° entradas (N)
- n° salidas (M)
- n° neuronas capa oculta (P_c)
- n° de capas ocultas (C). Una capa oculta es suficiente en la mayoría de los problemas, por lo q. trabajaremos con $C=1$ para simplificar

ALGORITMO DE ENTRENAMIENTO

El proceso de aprendizaje del perceptrón multicapa equivale a encontrar un mín. en la func. de error E :

$$E = \frac{1}{Q} \sum_{n=1}^Q e(n) \quad e(n) = \frac{1}{2} \sum_{i=1}^M (t_i(n) - y_i(n))^2$$

Para encontrar el mín. se recurre a técnicas de optimiz. no lineal, en las cuales la búsqueda sigue la dir. negativa del gradiente de E . En la práctica lo q. se minimiza son los errores de cada patrón, $e(n)$

- 1) Inicializ. pesos. ~~Heurística~~ valores aleatorios entre -0.5 y 0.5 → como inicio previo
- 2) Alimentación hacia adelante

$$z_{in j} = \sum_{i=0}^N x_i v_{ij} \quad z_j = f(z_{in j})$$

$$y_{in k} = \sum_{j=0}^P z_j w_{jk} \quad y_k = f(y_{in k})$$

Se puede definir una func. de activación diferente para cada neurona, pero en la práctica suele usarse la misma f por todas (sigmoide)

Para encontrar el mín. se recurre a técnicas de optimiz. no lineal, en las cuales la búsqueda sigue la dir. negativa del gradiente de E . En la práctica lo q. se minimiza son los errores de cada patrón, $e(n)$

- 1) Inicializ. pesos. ~~Heurística~~ valores aleatorios entre -0.5 y 0.5 → como inicio previo

4.2. Capa oculta

$$\delta_{in j} = \sum_{k=1}^M \delta_k w_{jk}$$

$$\delta_j = \delta_{in j} f'(z_{in j})$$

$$\Delta v_{ij} = \delta_j x_i$$

(crea delta generalizada)

5) Actualizar pesos

$$\Delta w_{jk} = \alpha \delta_k z_j$$

$$\Delta v_{ij} = \alpha \delta_j x_i$$

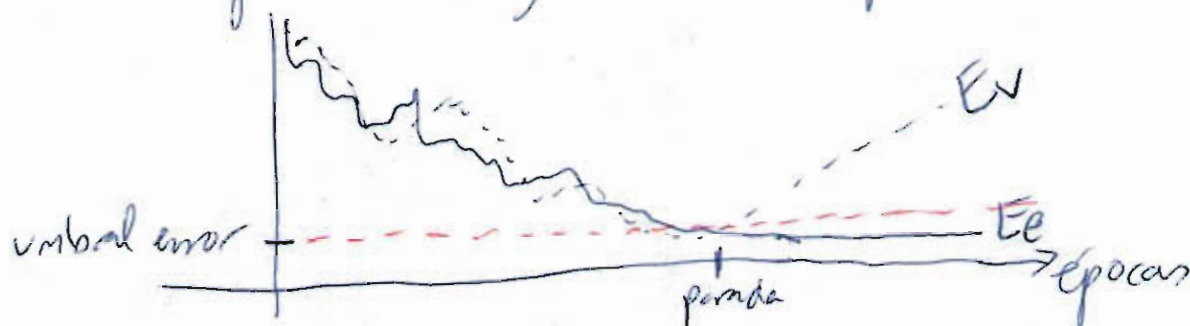
$$v_{ij}(t+1) = v_{ij}(t) + \Delta v_{ij}$$

$$w_{jk}(t+1) = w_{jk}(t) + \Delta w_{jk}$$

6) Evaluar cond de parada

Error entrenamiento

Se calcula el error E_e y se observa su evol. en cada época. En las primeras debe decrecer globalmente hasta sobrepasar un umbral preestablecido. Después se controla el error de validación E_v . Cuando comienza a crecer indica q. se está perdiendo generalidad y conviene parar



$$v_{ij}(t+1) = v_{ij}(t) + \Delta v_{ij}$$

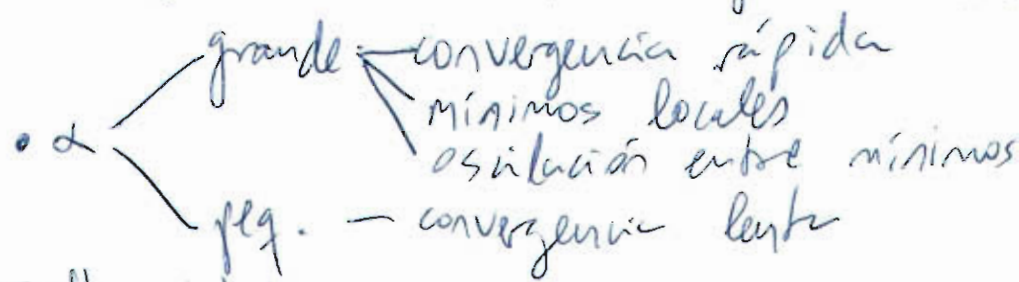
$$w_{jk}(t+1) = w_{jk}(t) + \Delta w_{jk}$$

6) Evaluar cond de parada

Error entrenamiento

Se calcula el error E_e y se observa su evol. en cada época. En las primeras debe decrecer globalmente hasta sobrepasar un umbral preestable-

El diseño de la red afecta a estos problemas:



• Pesos $\rightarrow$ Heurísticas $\rightarrow$ Nguyen-Widrow:

$$w_{ij} \in (-\beta, \beta), \quad \beta = 0.7 \sqrt{P}$$

• Codif. entradas y salidas: una entrada con valor 0 impide el aprendizaje, ralentizando la convergencia. Por eso se recomienda la codif. bipolar.

? • No es imprescindible normalizar entradas y salidas, pero obvio si no se normalizan las salidas hay q. utilizar la func. de activación identidad

• N° de parám. = $(N+1)P + P(M+1) = \boxed{P(N+M+1) + M}$

N y M suelen venir dados por el problema. Para no atascarse en mín. locales, conviene aumentar P; pero un P alto lleva al sobre-aprendizaje.

$\rightarrow$ Heurística: ~~aproximadamente~~ $10\% \cdot Q \leq \text{n° parám.} \leq 15\% \cdot Q$

También se puede intentar reducir N con técnicas

• Codif. entradas y salidas: una entrada con valor 0 impide el aprendizaje, ralentizando la convergencia. Por eso se recomienda la codif. bipolar.

? • No es imprescindible normalizar entradas y salidas, pero obvio si no se normalizan las salidas hay q. utilizar la func. de activación identidad

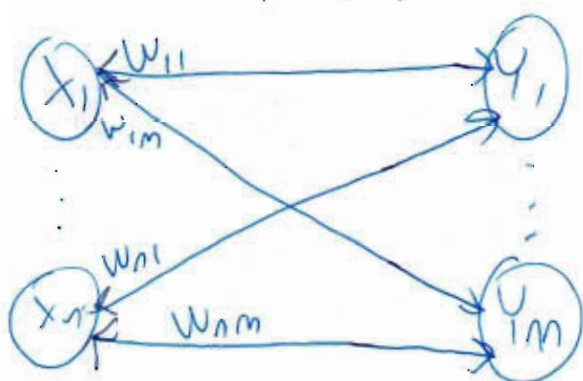
$$P(N+M+1) + M =$$

HETEROASOCIATIVA (de flujo directo)

TIPO: Asoc. patrones (ent. supervisado). pesos fijos

USOS:

ARQUITECTURA:



pesos simétricos

$m \neq n$ (Normal, $m < n$)

Codif. bipolar o binaria, pero la misma para entradas y salidas

ALGO. ENTREN. $w_{ij}(0) = 0$

No hay. los pesos se inicializan con la regla de Hebb o Delta. La matriz de pesos tb se puede calcular como la suma de las matrices y almacenar cada patrón por separado: $W_i = s(i) \cdot t(i)$; $W = \sum_{i=1}^p W_i$ ($p = n^\circ$ patrones)

ALGO. APLICACION

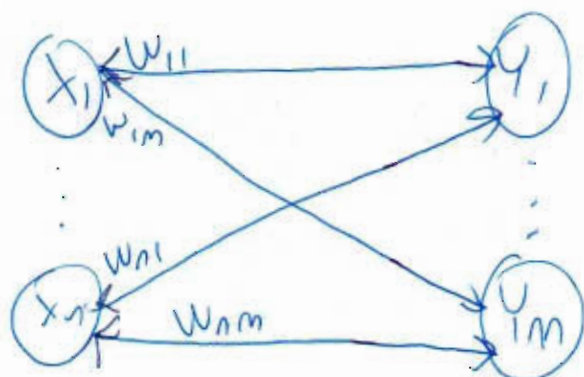
$$y_{-in j} = \sum_{i=1}^n x_i w_{ij}$$

Func. activación

y_j : _____

A partir de los pesos obtenidos durante el entrenamiento, se calcula la salida correspondiente para cada patrón

$$g(y) = \begin{cases} 1 & y > 0 \\ 0 & y = 0 \\ -1 & y < 0 \end{cases} \quad \text{bipolar}$$



pesos simétricos

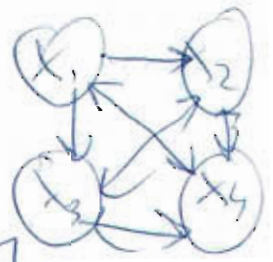
$m \neq n$ (Normal, $m < n$)

Codif. bipolar o binaria, pero la misma para entradas y salidas

ALGO. ENTREN. $w_{ij}(0) = 0$

AUTO ASOCIATIVA

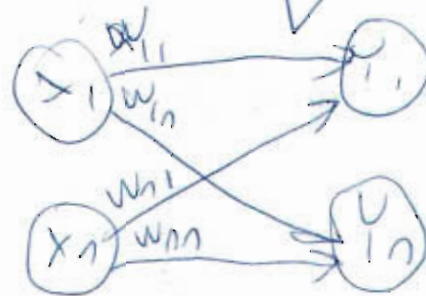
(~~lineal recurrente~~)
(no iterativa)



TIPO: Asoc. patrones (ent. supervisado)
Pesos fijos

USOS:

ARQUITECTURA:
≡ Heteroasociativa
Bipolar



ALGO. ENTREN.

Inic. pesos: $w_{ij} = 0$

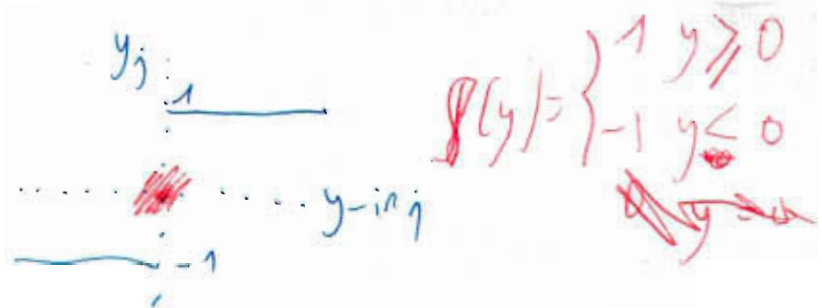
Regla aprendizaje: Hebb para vectores ortogonales ←

Poner a 0 los pesos de la diagonal (en Te, solo pueden almacenarse vectores ortogonales)

Normalmente, $w_{ii} = 0$ (sin auto-conexiones a red funciona mejor)

ALGO. APLICACIÓN

$$y_{-in j} = \sum_{i=1}^n x_i w_{ij}$$



≡ Heteroasociativa
Bipolar

ALGO. ENTREN.

Inic. pesos: $w_{ij} = 0$

Regla aprendizaje: Hebb para vectores ortogonales ←

Poner a 0 los pesos de la diagonal (en Te, solo pueden almacenarse vectores ortogonales)



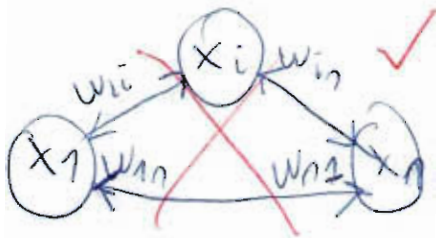
Normalmente, $w_{ii} = 0$ (sin auto-conexiones a red funciona mejor)

AUTOASOCIADOR CON FUNCIÓN UMbral

Tipo: Asoc. patrones cont. supervisado
Autoasociativa iterativa. Pesos fijos
(o local recurrente)

Usos:

ARQUITECTURA: como BSB pero sin autoconexiones
(todas conectados con todos)



$$w_{ii} = 0$$

Algo. ~~entrenamiento~~ ~~adaptación~~

Entradas y salidas binarias

FASE 1
(capacitación)
FASE 2
(aplicación)

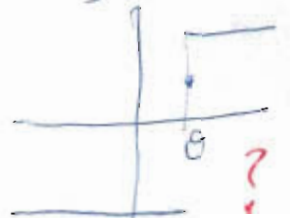
Inicializ. pesos (regla de Hebb). Normalmente, $\theta_i = 0$
(cada neurona tiene su propio umbral)

Activar la activación de todas las unidades:

$$x_{i\text{ new}} = \begin{cases} 1 & \text{si } \sum_j x_{j\text{ old}} w_{ij} > \theta_i \\ \text{cancelado} & \text{si } \sum_j x_{j\text{ old}} w_{ij} = \theta_i \\ -1 & \text{si } \sum_j x_{j\text{ old}} w_{ij} < \theta_i \end{cases}$$

$$x_{i\text{ new}} = \sum_j x_{j\text{ old}} w_{ij}$$

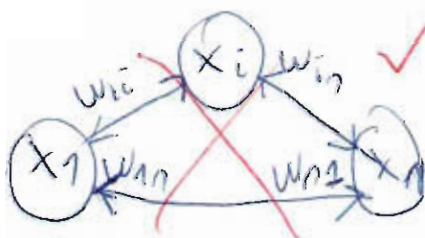
$$x = f(x_{i\text{ new}})$$



Condición de parada: x coincide con alguno
de los patrones almacenados en la memoria.

ARQUITECTURA:

como BSB pero sin autoconexiones
(todas conectados con todos)



$$w_{ii} = 0$$

Algo. ~~entrenamiento~~ ~~adaptación~~

Entradas y salidas binarias

FASE 1
(capacitación)

Inicializ. pesos (regla de Hebb). Normalmente, $\theta_i = 0$
(cada neurona tiene su propio umbral)

MIN-STATE-IN-A-BOX (BSB)

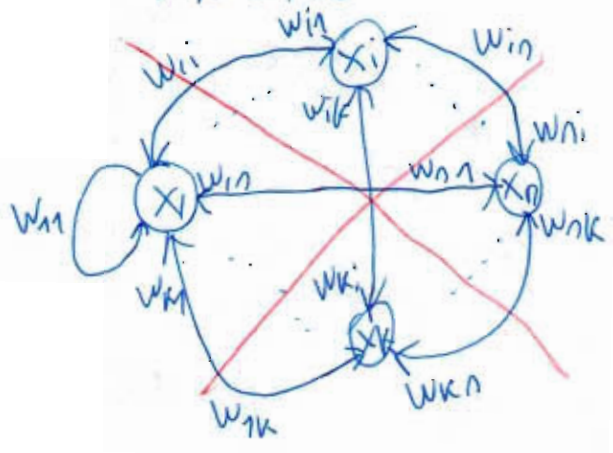
X

TIPO: Asociativa iterativa (Asoc. frases/ent. supervisada)
Recurrente

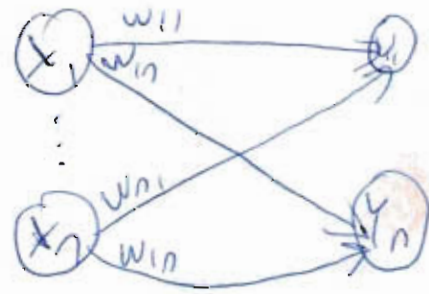
UJOS:

Bipolar

ARQUITECTURA



Matriz de pesos simétrica



ALGO. ENTREN.

Inic. pesos: valores aleatorios peg.
Tasas de aprendizaje α y β
Regla de aprendizaje: variante de Hebb
 $\Delta w_{ij} = \beta y_i y_j$

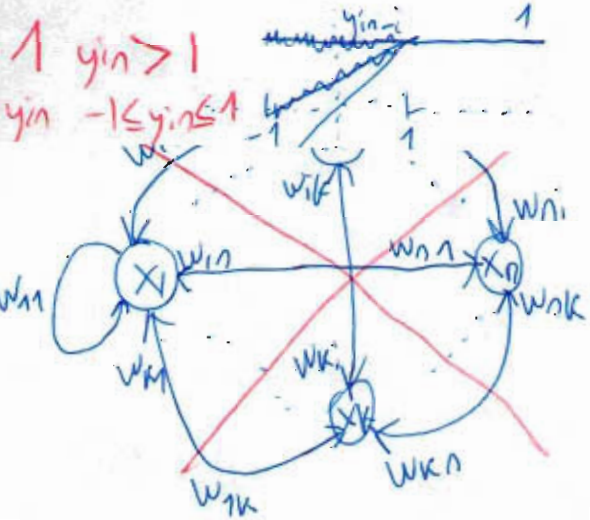
Func. activaci3n:

$y_i = x_i$

(codif. bipolar)

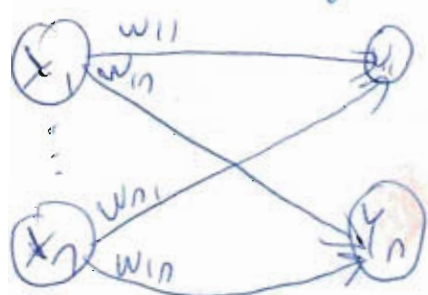
los valores se restringen a $(-1, 1)$

1 $y_i > 1$
 y_i $-1 \leq y_i \leq 1$



$$y_{ini} = y_i + \alpha \sum_{j=1}^n y_j w_{ji}$$

Matriz de pesos simétrica



ALGO. ENTREN.

... valores aleatorios peg.

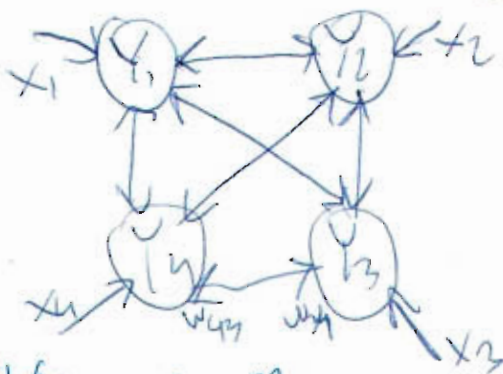
HOFFIELD DISCRETA

Bipolar o binario La matr. es siempre bipolar, así q. si los patrones son binarios hay q. convertirlos a bipolar antes de almacenarlos. **asoc. patrones/entran. supervisada**

TIPO: Autoasoc. iterativa recurrente. Pesos fijos

USOS: Mem. de acceso por contenido, reconocimiento de patrones (imágenes / voc...), optimización con restricciones (problema del viajante...)

ARQUITECTURA: = BSB + salvo pesos simétricos $w_{ii} = 0$



ALGO. ENTREN.

No tiene. Los pesos se calculan mediante la safe. fórmula (Hebb)

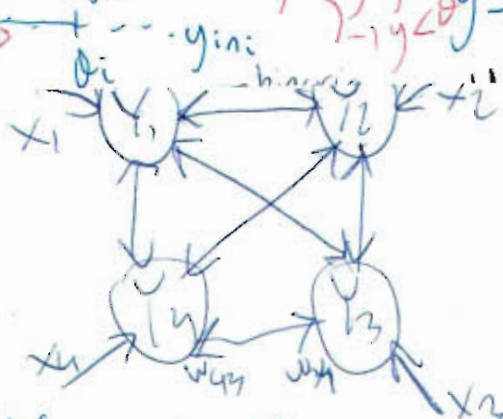
Hebb $w_{ij} = \begin{cases} 0 & \text{si } i=j \\ \sum_{p=1}^P s_i(p) s_j(p) \end{cases}$

simplemente los patrones en modo binario $\rightarrow w_{ij} = \sum_p [2s_i(p)-1][2s_j(p)-1]$

Hebb supervisada $\rightarrow s_i(p) t_j(p)$

ALGO. APLICACIÓN

$y = \begin{cases} 1 & y \geq 0 \\ 0 & y = 0 \\ -1 & y < 0 \end{cases}$



ALGO. ENTREN.

No tiene. Los pesos se

donde $s_i(p)$ son las componentes del patrón $s(p)$. P es el nº de patrones. Matricialmente,

$$\sum_p [S(p)^T S(p)]$$

suma cada matriz correspondiente al producto de $S(p)^T$ por $S(p)$ **escalar**

$$x_i + \sum_{j=1}^n y_j w_{ji}$$

La matr. de pesos siempre es bipolar, por lo si las entradas son binarias se convierten a bipolar durante la fase de entrenamiento. En la fase de aplic. las entradas se dejan en su codif. y el resultado se compara con los vectores alimpe. en la codif. orig.

Hebb $w_{ij} = \begin{cases} 0 & \text{si } i=j \end{cases}$

BIDIRECTIONAL ASSOCIATIVE MEMORY (BAM)

Tipo: ~~Autosoc.~~ iterativa (no supervisado, recurrente), pesos fijos

Uso: ~~heteroc.~~ x o y pueden ser incompletos / erróneos

→ Binaria o bipolar (mejor bipolar) → binario se convierte a bipolar

Algoritmo

La recuperación es iterativa: mientras no se converge, se calculan alternativamente X e Y copiando la bidireccionalidad

Pesos sinápticos $X \rightarrow Y = W$
(Simétricos) $Y \rightarrow X = W^T$

Peso $x_i \rightarrow y_j = \text{Peso } y_j \rightarrow x_i$

11	12
21	22
31	32
41	42

Hdb patrones binarios

$$W = \sum_p (2s(p)^T - 1) \cdot (2t(p) - 1)$$

Algo. ENTREN. / APLICACIÓN

FASE 1 Inic. pesos: Regla Hebb (patrones bipolares) $\downarrow$ entrada p objetivo p

(Calcularlos)

$$w_{ij} = \sum_p s_i(p) t_j(p)$$

$$s_i(p), t_i(p)$$

FASE 2 Presentar patrones a la capa X X, Y

Func. activación

← un patrón incompleto en X o en Y sirve de pto. de partida, ayudando a completar posiciones y la red por sí sola lo puede

Calcular respuesta capa Y (y_i)

$$y_{in i} = \sum_j w_{ij} x_j \quad y_i = f(y_{in i})$$

Calcular respuesta capa X (x_i)

$$x_{in i} = \sum_j w_{ji} y_j \quad x_i = f(x_{in i})$$

(Simétricos)

$$Y \rightarrow X = W^T$$

Peso $x_i \rightarrow y_j = \text{Peso } y_j \rightarrow x_i$

11	12
21	22
31	32
41	42

Hdb patrones binarios

$$W = \sum_p (2s(p)^T - 1) \cdot (2t(p) - 1)$$

Algo. ENTREN. / APLICACIÓN

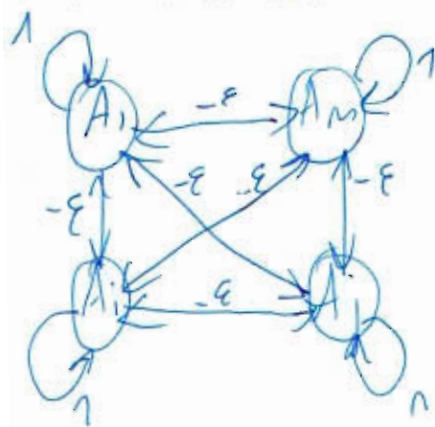
FASE 1 Inic. pesos: Regla Hebb (patrones bipolares) $\downarrow$ entrada p objetivo p

MAXNET

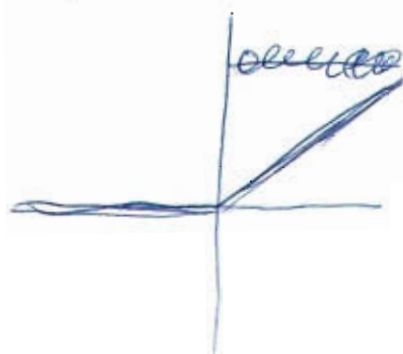
TIPO: Costo. continua positiva
competitiva pero fijo

USOS:

ARQUITECTURA



FUNC. ACTIVACION



$x > 0$
 $x \leq 0$

ALGO. ENTREN.

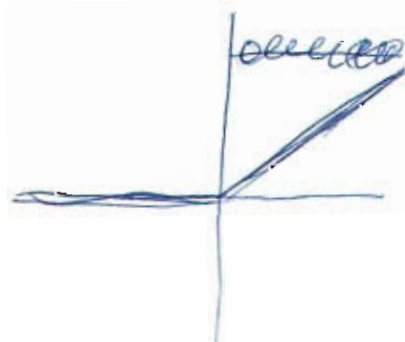
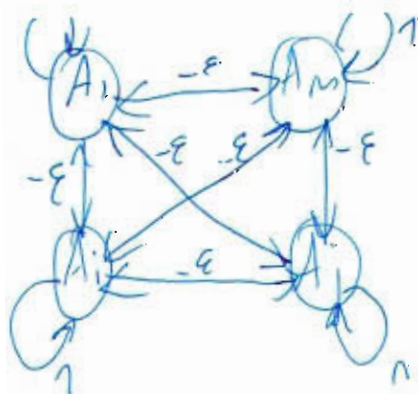
No tiene. $w_{ij} = \begin{cases} 1 & \text{si } i=j \\ -\epsilon & \text{si } i \neq j \end{cases} \quad (0 < \epsilon < \frac{1}{m})$

ALGO. APLICACIÓN

Actualizar activaciones:

$$a_i(\text{new}) = [a_i(\text{old}) - \epsilon \sum_{k \neq i} a_k(\text{old})] = \max(0, \sum w_{ij} a_j(\text{old}))$$

Cond. parada: sólo hay uno nodo activado



$x > 0$
 $x \leq 0$

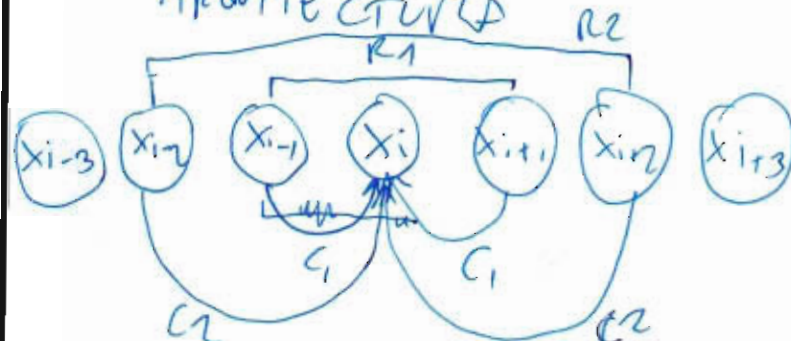
ALGO. ENTREN.

SOMbrero MEXICANO

Tipo: competitiva pero fijo

Uso: refuerzo contraste

ARQUITECTURA



R_1 : conexiones excitadoras
 R_2 : conexiones inhibidoras

Codif. continua

R_1 : radio de cooperación refuerzo positivo
 $(W_K = C_1, C_1 > 0)$

R_2 : radio competición
 $(W_K = C_2, C_2 < 0)$

región cooperación: $\{x_i, x_{i+1}, x_{i-1}\}$
 competición: $\{x_{i-2}, x_{i+2}\}$

ALGO. APLICACIÓN

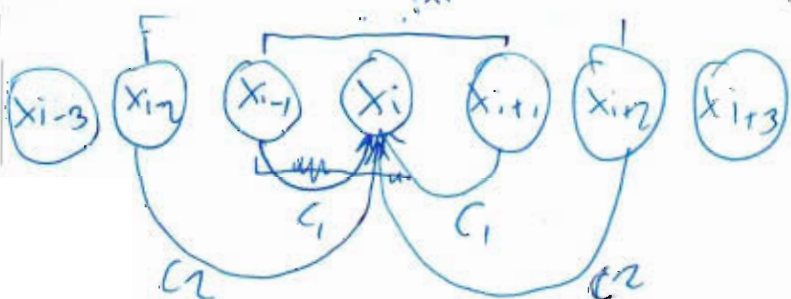
Para cada unidad de activación x_i , calcular la entrada red,

$$x_{ini} = C_1 \sum_{k=-R_1}^{R_1} x_{i+k}^{(old)} + C_2 \sum_{k=-R_2}^{-R_1-1} x_{i+k}^{(old)} + C_2 \sum_{k=R_1+1}^{R_2} x_{i+k}^{(old)}$$

región cooperación
 (anillo interno)

región competición (anillo externo)
 $x_i^{(old)} = x_i$

Func. activación



R_1 : conexiones excitadoras
 R_2 : conexiones inhibidoras

$$x_i = f(x_{ini})$$

R_1 : radio de refuerzo positivo
 $(W_K = C_1, C_1 > 0)$

R_2 : radio competición
 $(W_K = C_2, C_2 < 0)$

región cooperación: $\{x_i, x_{i+1}, x_{i-1}\}$
 competición: $\{x_{i-2}, x_{i+2}\}$

HAMMING

Bipolar / binaria

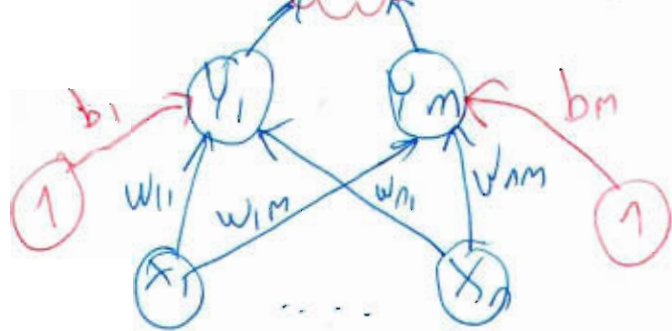
no tiene sentido codif. binaria

TIPO: Competitivo peso fijo

USOS: determinar a q. vector ejemplo se parece + el vector de entrada

ARQUITECTURA

MAXNET → elige la unidad y + activa



Algo. ~~Algoritmo~~
 Inic. pesos: $w_{ij} = \frac{e_i \cdot e_j}{2}$
 bias: $b_j = \frac{n}{2}$

$e(j)$ es el patrón de referencia (ejemplo) de la unidad Y_j
 $e_i(j)$ es cada una de sus componentes

Para cada patrón de entrada,

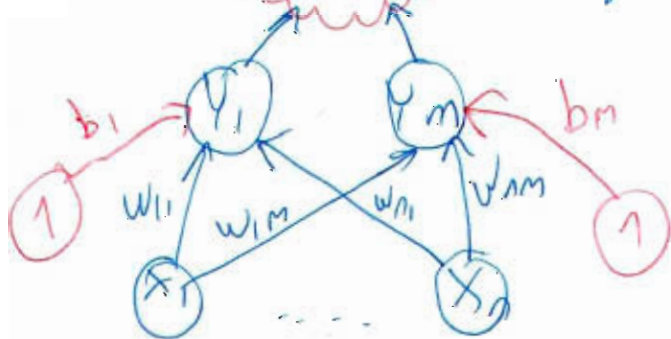
entrada net: $y_{in j} = b_j + \sum_{i=1}^n x_i \cdot w_{ij}$

(x e y son vectores bipolares)

$y_j = \max_j y_{in j}$

Buscar el mejor ejemplo con MAXNET aplicada sobre el vector y

MAXNET → elige la unidad y + activa



carrito de supermercado

Algo. ~~Algoritmo~~

MAPA AUTOORGANIZATIVO (KOHONEN)

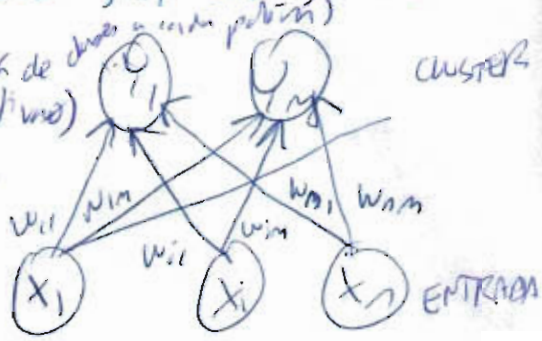
cualquier codif.

TIPO: Competitiva / entren. no supervisada, recurrente

USOS: agrupar ^(clasificar) los patrones en clusters sin conocimiento

- 2) q. la dist. se establece (no hay cambios en la asignación de datos a cada patrón)
- 3) q. la dist. oscila (cada patrón se asigna a una clase alternativa)

ARQUITECTURA: = heteroasociativa (pero con topología)



El n° de datos (salidas) se desconoce a priori, Alto. ENTREN. por lo q. se estima

Inicializ. pesos: pueden reflejar un conocimiento previo, o simplemente ser valores aleatorios

Elegir radio estándar R y vecindad (lineal, rectangular, hexagonal)²

Inic. tasa aprendizaje α y tasa de variación de R

para cada patrón: Buscar la neurona j + próxima al vector de entrada x [aquella con la distancia D menor, $D(j) = \sum_{i=1}^n (w_{ij} - x_i)^2$]

Regla aprendizaje ~~Kohonen~~: $\Delta w_{ij} = \alpha (x_i - w_{ij}(t))$ respecto a los pesos

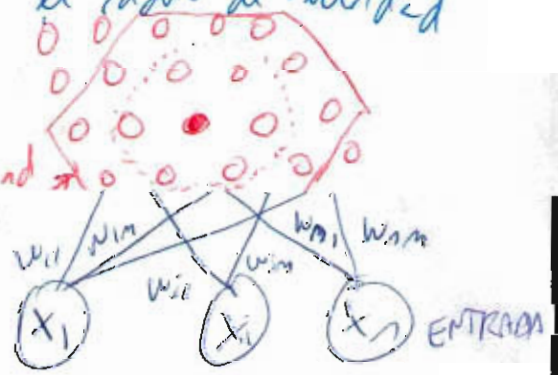
Actualizar pesos de la unid. ganadora y_j y sus vecinos

En cada época se actualiza α y el radio de vecindad

Cond. parada: iteraciones $\geq 500 \cdot n$

COMENTARIOS algunas

ARQUITECTURA: = heteroasociativa (pero con topología)



El n° de datos (salidas) se desconoce a priori, Alto. ENTREN. por lo q. se estima

Inicializ. pesos: pueden reflejar un conocimiento previo, o simplemente ser valores aleatorios

Elegir radio estándar R y vecindad (lineal, rectangular, hexagonal)²

Inic. tasa aprendizaje α y tasa de variación de R

LEARNING VECTOR QUANTIZATION (LVQ)

Unsupervised
code

TIPO: Asoc. patrones / entren. ^{#2} Supervisado, FF, competitiva

USOS: ^{Clasificar patrones nuevos en base a otros conocidos}
SOM para aprendizaje supervisado

ARQUITECTURA

= ~~eq.~~ SOM, pero sin topología. Se conoce a q. clase pertenece cada patrón (es decir, se sabe de antemano el n.º de clases, y por tanto, el n.º de salidas)

ALGO. ENTREN. (LVQ1)

Inicializ. vectores de pesos de referencia. Opciones:

1/ Usar los m 1.º vectores de entrenamiento

2/ Inicializ. los pesos aleatoriamente

Inicial. tasa aprendizaje α

Buscar el ~~vector~~ ^{la una de salida} con menor distancia

^{euclídea} $\rightarrow ||x - w_j|| = \sqrt{\sum (x_i - w_{ij})^2}$
entre el patrón y los pesos de ref.

Si la categoría correcta del patrón es = a la categoría obtenida $\rightarrow$

$w_j(\text{new}) = w_j(\text{old}) + \Delta w_j$

Lo contrario $\rightarrow w_j(\text{new}) = w_j(\text{old}) - \Delta w_j$

$\Delta w_j = \alpha (x - w_j(\text{old}))$

(Kohonen)

$\rightarrow$ Reducir α (en cada época)

Cond. parada — n.º de iteraciones

= ~~eq.~~ SOM, pero sin topología. Se conoce a q.

clase pertenece cada patrón (es decir, se sabe de antemano el n.º de clases, y por tanto, el n.º de salidas)

ALGO. ENTREN. (LVQ1)

Inicializ. vectores de pesos de referencia. Opciones:

1/ Usar los m 1.º vectores de entrenamiento

La raíz cuadrada no es necesaria
 $\rightarrow ||x - w_j|| = \sqrt{\sum (x_i - w_{ij})^2}$

CONTRAPROPAGACION

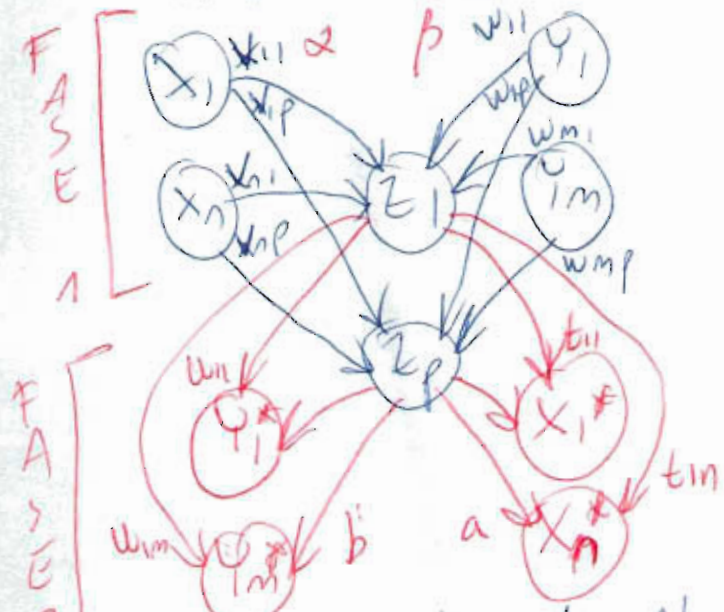
(Counterpropagation)

Codif. Continua Real

TIPO: Competitiva

Usos: Aprox. func. } contraprop. total
 Asociar patrones }
 Comparar info. (contrapropagacion de flujo directo)

ARQUITECTURA



Contraprop. total / completa

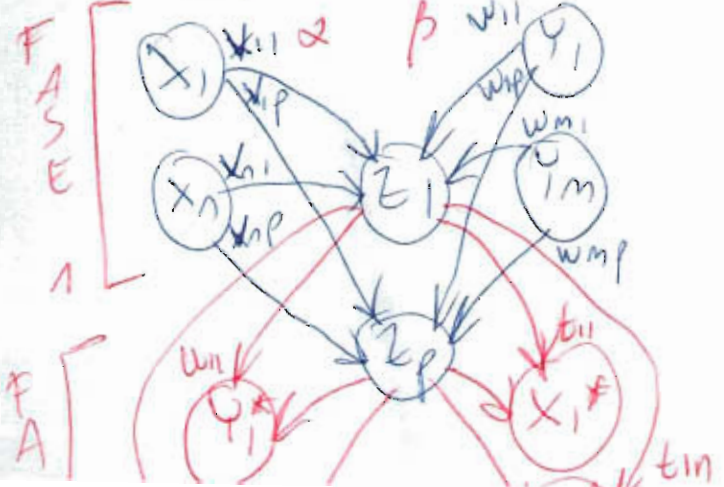
Func. invertible: $x \rightarrow y$

X, Y : entradas; Z : cluster; X^*, Y^* : salidas

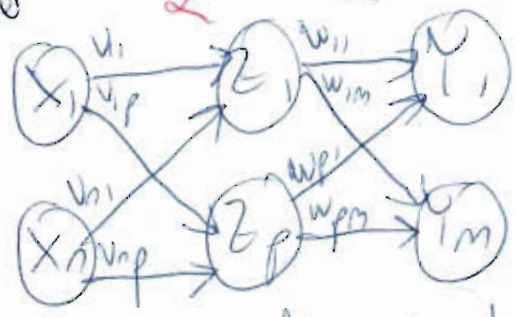
ALGO. ENTREN.

Inicializar pesos: v, w, t, u

ARQUITECTURA

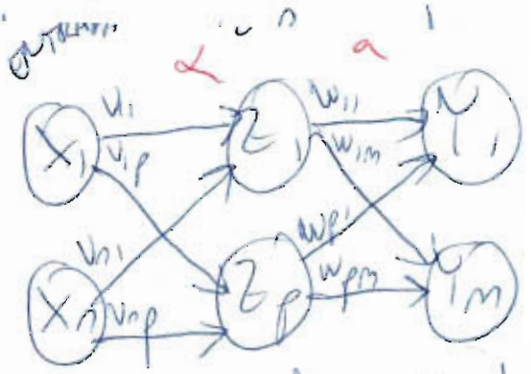


ENTRADA CLUSTER SALIDA



Contraprop. flujo directo / sólo hacia adelante

Func. no invertible: $x \rightarrow y$
 La red de flujo directo se utiliza = q. la total, salvo q. obvio no hay pesos u y t , por lo q. en la fase 1 se calculan sólo los v y w en la 2 los t (q. aquí se etiquetan como w y t)



Contraprop. flujo directo / sólo hacia adelante

Func. no invertible: $x \rightarrow y$

Actualizar pesos v y w con Δv y Δw
 Reducir factores aprendizaje α y β (después de cada época, no de cada patrón)
 Cond. parada Fase 1 (pesos no cambian, entradas asociadas siempre a la misma salida, factor aprendizaje muy pequeño)
 Fase 2: refinamiento pesos cluster $\rightarrow$ salida
 (en esta fase, α y β han quedado reducidos a valores pequeños)
 Entradas $\rightarrow X = x, Y = y$
 Buscar unid. ganadora J
 Ajuste pesos $\rightarrow v, w \rightarrow \min$ (α y β son muy pequeños)
 $\Delta u_{jk} = a (y_k - u_{jk}^{old})$
 $\Delta t_{ji} = b (x_i - t_{ji}^{old})$
 por lo que puede considerarse como delta $\leftarrow$ objetivos de x_k e y_i $\rightarrow$ activación unids. x_k e y_i
 Reducir factores aprendizaje a y b
 Cond. parada Fase 2

APLICACIÓN

Suministrar patrones x e y a la red (x o y pueden ser desconocidos), con lo q. la red activa una neurona Z_j del cluster. Los pesos u_{jk} y t_{ji} asociados a dicha neurona son la $u_{jx_1}, u_{jx_2}, \dots, u_{jx_n}$ y $t_{jy_1}, t_{jy_2}, \dots, t_{jy_n}$.
 Buscar unid. ganadora J
 Ajuste pesos $\rightarrow v, w \rightarrow \min$ (α y β son muy pequeños)
 $\Delta u_{jk} = a (y_k - u_{jk}^{old})$
 $\Delta t_{ji} = b (x_i - t_{ji}^{old})$
 por lo que puede considerarse como delta $\leftarrow$ objetivos de x_k e y_i $\rightarrow$ activación unids. x_k e y_i
 Reducir factores aprendizaje a y b

ART

Adaptive Resonance Theory

Diseñada para ser estable y plástica

Tipo: No supervisado

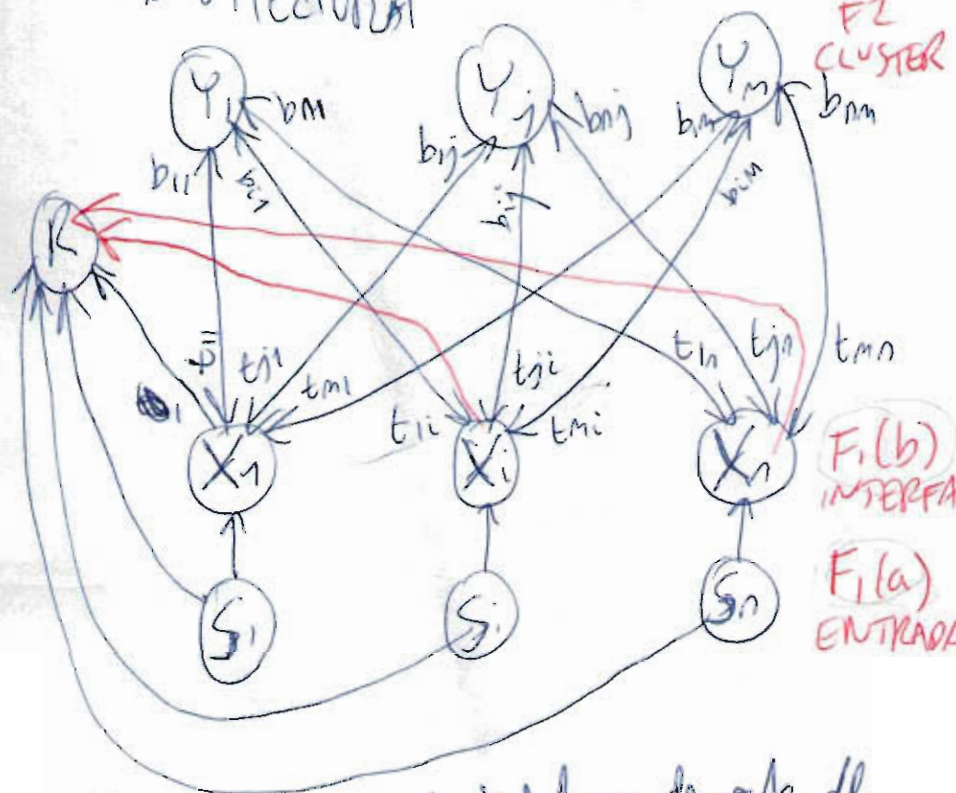
ART 1: entradas binarias

ART 2: entradas continuas

binaria, bipolar?

- Con el entrenamiento se consigue q. en los pesos tsi se almacenen a los representantes de cada categoría
- Se hace entre la info. sin y la info. procesada de la capa competitiva

Usos: Clasificación de patrones q. llegan al campo geniculado y la corteza visual primaria
ART 2: HECUNA



n = nº componentes vector entrada
 m = nº máx. de clusters
 b_{ij} = pesos $F_1(b) \rightarrow F_2$
 t_{ji} = pesos $F_2 \rightarrow F_1$
Las neuronas se representan en mayúscula, y su activación en minúscula: Y_j vs y_j
 $R \rightarrow$ mecanismo de Reset (previo a la actualiz. de los pesos)

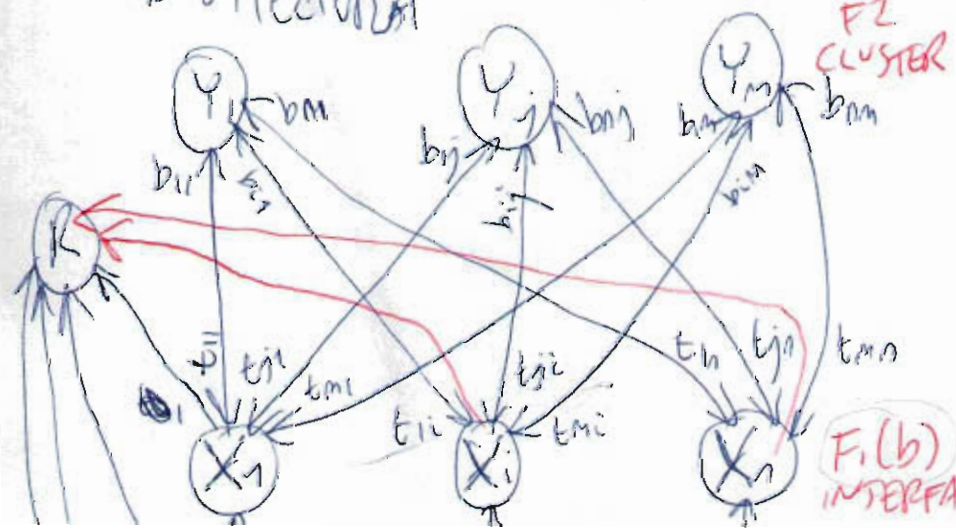
- en esta red no puede hablarse de regla de aprendizaje, es una simple actualiz. de pesos

ALGO. ENTRENAMIENTO (ART1)

Inicializar parámetros: $L > 1$

Inicializar pesos: $0 < b_{ii}(0) < 1$ / L : $t_{ii}(0) = 1$

grado de similitud entre los patrones de un mismo cluster
factor de vigilancia
 $0 < \rho \leq 1$



n = nº componentes vector entrada
 m = nº máx. de clusters
 b_{ij} = pesos $F_1(b) \rightarrow F_2$
 t_{ji} = pesos $F_2 \rightarrow F_1$
Las neuronas se representan en mayúscula, y su activación en minúscula: Y_j vs y_j

- Calcular la norma de s^1 (el vector de entrada):

$$\|s\| = \sum_i s_i$$

- Mandar la señal de entrada de $F_i(a)$ a $F_i(b)$

$$x_i = s_i$$

- Para cada nodo j de F_2 (q. no está inhibido, $y_j \neq -1$), calcular el valor de activación $y_j = \sum_i b_{ij} x_i \rightarrow y = x \cdot b$ (vectorial)

- Mientras Reset sea verdadero,

- encontrar la neurona J^5 ganadora "laquelle tal q."

$$y_J \geq y_j \quad \forall j$$

- recalcular las activaciones de $F_1(b)$

$$x_i = s_i t_{ji} \quad x = s t_J \quad (\text{escalar})$$

- calcular la norma de x

$$\text{Test de Reset } \|x\| = \sum_i x_i$$

Si $\frac{\|x\|}{\|s\|} < \rho \rightarrow y_J = -1$ (inhibir nodo ganador) (divinizado de la competición)

Si no, actualizar pesos para nodo J^5 y poner Reset = False:

$$b_{iJ}(\text{new}) = \frac{L x_i}{L - 1 + \|x\|}$$

$$t_{ji}(\text{new}) = x_i \quad (\text{una vez q. se asigna 0 a } t_{ji}, \text{ ya no puede volver a valer 1, con lo q. se consigue la estabilidad})$$

- Evaluar condición de parada (una de las 3 sigtes):

- los pesos no han cambiado

- no se han reiniciado unids.

- se ha alcanzado el n.º máx. de épocas

el valor de activación

- Mientras Reset sea verdadero,

- encontrar la neurona J^5 ganadora "laquelle tal q."

$$y_J \geq y_j \quad \forall j$$

- recalcular las activaciones de $F_1(b)$

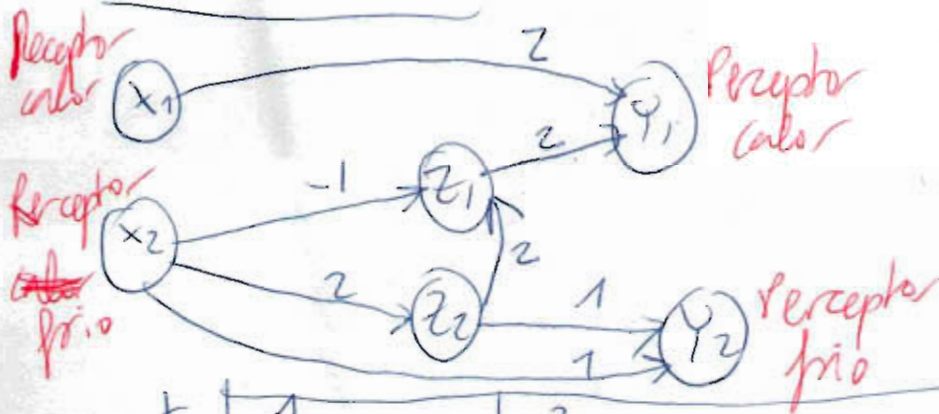
$$x_i = s_i t_{ji} \quad x = s t_J \quad (\text{escalar})$$

- calcular la norma de x

$$\|x\| = \sum_i x_i$$

EJEMPLOS

McCulloch - Pitts



Todas las neuronas tienen un umbral de activación 2

t	1	2
x ₁	0 0 1 1	0 0 0 0 1 1 1 0 0 1
x ₂	0 1 0 1	1 1 1 0 0 0 0 1 1 0
z ₁	0 -1 0 -1	-1+4 3 4 0+4 0 0+0 -1+0 0+4
z ₂	0 2 0 2	2 0 0 2 0
y ₁	0 0 0 2	0+2 -2 2+0 0+0 2+-2
y ₂	0 1 2 1	1+2 2 0+0 1+0 0+2

frio calor frio frio calor - frio t=2

t=1 x₁=1 (calor) y₂=1 (frio)
x₂=1 (frio) y₂=1 (frio)
x₁=1 (calor) y₁=1 (calor)
x₂=1 (frio) y₁=y₂=0
x₁=1 (calor) y₂=1 (frio)
x₂=1 (frio) y₂=1 (frio)

Observando los valores se deduce q. si se recibe frio en 2 instantes consecutivos, se percibe frio en ambos. Si se recibe calor y luego en los 2 instantes, 1º se percibe frio y luego calor. Si se recibe

frio

t	1	2
x ₁	0 0 1 1	0 0 0 0 1 1 1 0 0 1
x ₂	0 1 0 1	1 1 1 0 0 0 0 1 1 0
z ₁	0 -1 0 -1	-1+4 3 4 0+4 0 0+0 -1+0 0+4
z ₂	0 2 0 2	2 0 0 2 0
y ₁	0 0 0 2	0+2 -2 2+0 0+0 2+-2
y ₂	0 1 2 1	1+2 2 0+0 1+0 0+2

calor frio t=2

Intro. al tiempo

$$Y_1(t) = x_1(t-1) \text{ OR } z_1(t-1)$$

$$Y_2(t) = x_2(t-1) \text{ AND } z_2(t-1)$$

$$z_1(t) = z_2(t-1) \text{ AND NOT } x_2(t-1)$$

$$\hookrightarrow z_1(t-1) = z_2(t-2) \text{ AND NOT } x_2(t-2)$$

$$z_2(t) = x_2(t-1) \rightarrow z_2(t-2) = x_2(t-3)$$

$$Y_1(t) = x_1(t-1) \text{ OR } (z_2(t-2) \text{ AND NOT } x_2(t-2))$$

$$= \boxed{x_1(t-1) \text{ OR } (x_2(t-3) \text{ AND NOT } x_2(t-2))}$$

$$Y_2(t) = \boxed{x_2(t-1) \text{ AND } x_2(t-2)}$$

Es decir, se percibe calor si: se ha ^{recibido} ~~percebido~~ calor en el instante anterior, o se ha ^{recibido} ~~experimentado~~ frío solo un instante. Se percibe frío si durante 2 instantes se ha ^{recibido} ~~percebido~~ frío

$$Y_1(t) = x_1(t-1) \text{ OR } (z_2(t-2) \text{ AND NOT } x_2(t-2))$$

$$= \boxed{x_1(t-1) \text{ OR } (x_2(t-3) \text{ AND NOT } x_2(t-2))}$$

$$Y_2(t) = \boxed{x_2(t-1) \text{ AND } x_2(t-2)}$$

Es decir, se percibe calor si: se ha ^{recibido} ~~percebido~~ calor en el instante anterior, o se ha ^{recibido} ~~experimentado~~ frío solo un instante. Se percibe frío si durante 2 instantes se ha ^{recibido} ~~percebido~~ frío

Hebb

Func. AND

$$x(1) = (1 \quad 1)$$

$$x(2) = (1 \quad -1)$$

$$x(3) = (-1 \quad 1)$$

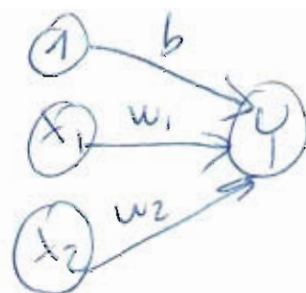
$$x(4) = (-1 \quad -1)$$

$$t(1) = 1$$

$$t(2) = -1$$

$$t(3) = -1$$

$$t(4) = -1$$



Inicializ. pesos

$$w = (0 \quad 0)$$

$$x(1) \rightarrow w_i(\text{new}) = w_i(\text{old}) + x_i t_i$$

$$w_1(\text{new}) = 0 + 1 \cdot 1 = 1$$

$$w_2(\text{new}) = 0 + 1 \cdot 1 = 1$$

$$b(\text{new}) = 0 + 1 \cdot 1 = 1$$

$$\left. \begin{array}{l} 1 + x_1 + x_2 = 0 \\ \downarrow \\ x_2 = -x_1 - 1 \end{array} \right\} R1$$

$$x(2) \rightarrow w_1(\text{new}) = 1 + (1 \cdot (-1)) = 0$$

$$w_2(\text{new}) = 1 + (-1 \cdot -1) = 2$$

$$b(\text{new}) = 1 - 1 = 0$$

$$\left. \begin{array}{l} 2x_2 = 0 \\ \downarrow \\ x_2 = 0 \end{array} \right\} R2$$

$$x(3) \rightarrow w_1(\text{new}) = 0 + (-1 \cdot -1) = 1$$

$$w_2(\text{new}) = 2 + (-1 \cdot -1) = 3$$

$$b(\text{new}) = 0 - 1 = -1$$

$$\left. \begin{array}{l} -1 + x_1 + x_2 = 0 \\ \downarrow \\ x_2 = 1 - x_1 \end{array} \right\} R3$$

$$x(4) \rightarrow w_1(\text{new}) = 1 + (-1 \cdot -1) = 2$$

$$w_2(\text{new}) = 3 + (-1 \cdot -1) = 4$$

$$b(\text{new}) = -1 + -1 = -2$$

$$\left. \begin{array}{l} -2 + 2x_1 + 4x_2 = 0 \\ \downarrow \\ x_2 = \frac{2 - 2x_1}{2} = 1 - x_1 \end{array} \right\} R4$$

$$x(3) - x_2 + x(1)$$

Inicializ. pesos

$$w = (0 \quad 0)$$

$$x(1) \rightarrow w_i(\text{new}) = w_i(\text{old}) + x_i t_i$$

$$w_1(\text{new}) = 0 + 1 \cdot 1 = 1$$

$$w_2(\text{new}) = 0 + 1 \cdot 1 = 1$$

$$b(\text{new}) = 0 + 1 \cdot 1 = 1$$

$$\left. \begin{array}{l} 1 + x_1 + x_2 = 0 \\ \downarrow \\ x_2 = -x_1 - 1 \end{array} \right\} R1$$

$$x(2) \rightarrow w_1(\text{new}) = 1 + (1 \cdot (-1)) = 0$$

$$w_2(\text{new}) = 1 + (-1 \cdot -1) = 2$$

$$\left. \begin{array}{l} 2x_2 = 0 \\ \downarrow \\ x_2 = 0 \end{array} \right\} R2$$

Perceptron

func. AND

$$x(1) = (0 \quad 0)$$

$$x(2) = (0 \quad 1)$$

$$x(3) = (1 \quad 0)$$

$$x(4) = (1 \quad 1)$$

$$t(1) = -1$$

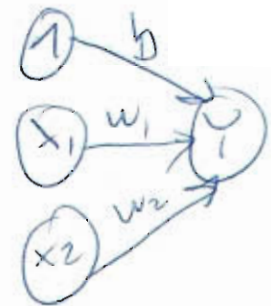
$$t(2) = -1$$

$$t(3) = -1$$

$$t(4) = 1$$

$$\alpha = 1$$

$$\theta = 0.2$$



$$w_1 = w_2 = b = 0$$

$$x(1) \quad y_{in} = 0 \rightarrow y = 0 \quad y \neq t(1)$$

$$w_i(new) = w_i(old) + \alpha t x_i$$

$$w_1(new) = 0 + 1 \cdot (-1) \cdot 0 = 0$$

$$w_2(new) = 0 + 1 \cdot (-1) \cdot 0 = 0$$

$$b(new) = 0 + 1 \cdot (-1) = -1$$

$$x(2) \quad y_{in} = -1 \rightarrow y = -1 \quad y \neq t(2)$$

$$w_1(new) = 0 + 1 \cdot (-1) \cdot 0 = 0$$

$$w_2(new) = 0 + 1 \cdot (-1) \cdot 1 = -1$$

$$b = -1 + 1 \cdot (-1) = -2$$

no define recta

$$-2 - x_2 = \pm \theta = \pm 0.2$$

$$x_2 = -2 \pm 0.2 \quad R1$$

$$x_2' = -1.8 \quad R1$$

$$x(3) \quad y_{in} = -1 \rightarrow y = -1 \quad y = t(3)$$

$$x(4) \quad y_{in} = -1 \rightarrow y = -1 \quad y \neq t(4)$$

$$w_1(new) = 0 + 1 \cdot 1 = 1$$

$$w_2(new) = 0 + 1 \cdot 1 = 1$$

$$b(new) = -1 + 1 = 0$$

$$x_1 + x_2 = \pm 0.2$$

$$x_2 = 0.2 - x_1$$

$$x_2' = -0.2 - x_1 \quad R2$$

$$R2' \quad R2 \quad x(3) \quad x_2' \quad R3$$

Tras esta 1ª época, los pesos son
(0 1 1)

$$w_1 = w_2 = b = 0$$

$$x(1) \quad y_{in} = 0 \rightarrow y = 0 \quad y \neq t(1)$$

$$w_i(new) = w_i(old) + \alpha t x_i$$

$$w_1(new) = 0 + 1 \cdot (-1) \cdot 0 = 0$$

$$w_2(new) = 0 + 1 \cdot (-1) \cdot 0 = 0$$

$$b(new) = 0 + 1 \cdot (-1) = -1$$

$$x(2) \quad y_{in} = -1 \rightarrow y = -1 \quad y \neq t(2)$$

$$w_1(new) = 0 + 1 \cdot (-1) \cdot 0 = 0$$

$$w_2(new) = 0 + 1 \cdot (-1) \cdot 1 = -1$$

no define recta

$$-2 - x_2 = \pm \theta = \pm 0.2$$

$$x_2 = -2 \pm 0.2 \quad R1$$

Adaline

Func. OR

$$x(1) = (1 \quad 1)$$

$$x(2) = (1 \quad -1)$$

$$x(3) = (-1 \quad 1)$$

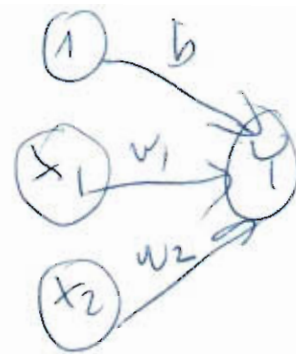
$$x(4) = (-1 \quad -1)$$

$$t(1) = 1$$

$$t(2) = 1$$

$$t(3) = 1$$

$$t(4) = -1$$



$$\alpha = 0.4$$

Init. pesos: $(\overset{b}{0.2} \quad \overset{w_1}{0.1} \quad \overset{w_2}{0.2})$

$$x(1) \rightarrow y_{in} = 0.2 + 0.1 + 0.2 = 0.5$$

$$w_{ij}(\text{new}) = w_{ij}(\text{old}) + \alpha (t_j - y_{in,j}) x_i$$

$$w_1(\text{new}) = 0.1 + 0.4(0.5 - 1) \cdot 1 = 0.3$$

$$w_2(\text{new}) = 0.2 + 0.4(0.5 - 1) \cdot 1 = 0.4$$

$$b(\text{new}) = 0.2 + 0.4(0.5 - 1) = 0.4$$

$$x(2) \rightarrow y_{in} = 0.4 + 0.3 - 0.4 = 0.3$$

$$w_1(\text{new}) = 0.3 + 0.4(0.3 - 1) \cdot 1 = 0.58$$

$$w_2(\text{new}) = 0.4 + 0.4(0.3 - 1) \cdot (-1) = 0.12$$

$$b(\text{new}) = 0.4 + 0.28 = 0.68$$

$$x(3) \rightarrow y_{in} = 0.68 - 0.58 + 0.12 = 0.22$$

$$w_1(\text{new}) = 0.58 + 0.4(0.22 - 1) \cdot (-1) = 0.27$$

$$w_2(\text{new}) = 0.12 + 0.4(0.22 - 1) \cdot 1 = 0.43$$

$$b(\text{new}) = 0.68 + 0.4(0.22 - 1) = 0.99$$

$$x(4) \rightarrow y_{in} = 0.99 - 0.27 - 0.43 = 0.29$$

$$w_1(\text{new}) = 0.27 + 0.4(0.29 - 1) \cdot (-1) = 0.78$$

Init. pesos: $(\overset{b}{0.2} \quad \overset{w_1}{0.1} \quad \overset{w_2}{0.2})$

$$x(1) \rightarrow y_{in} = 0.2 + 0.1 + 0.2 = 0.5$$

$$w_{ij}(\text{new}) = w_{ij}(\text{old}) + \alpha (t_j - y_{in,j}) x_i$$

$$w_1(\text{new}) = 0.1 + 0.4(0.5 - 1) \cdot 1 = 0.3$$

$$w_2(\text{new}) = 0.2 + 0.4(0.5 - 1) \cdot 1 = 0.4$$

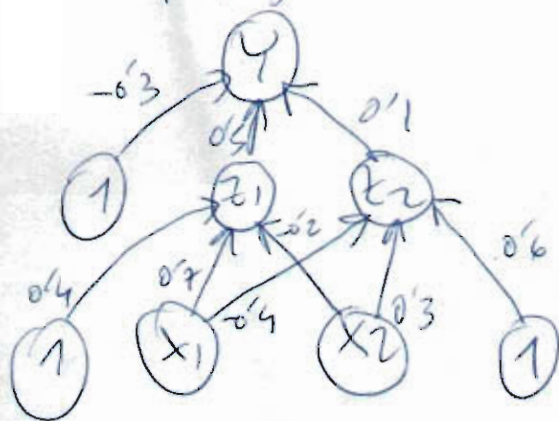
$$b(\text{new}) = 0.2 + 0.4(0.5 - 1) = 0.4$$

$$x(2) \rightarrow y_{in} = 0.4 + 0.3 - 0.4 = 0.3$$

Para saber el valor correspondiente a $(-1 \ 1)$,
usamos la red con la func. de activación escalón.

$$y_{in} = 0's - 0's + 0's = 0's \rightarrow y = 1$$

Retropropagación



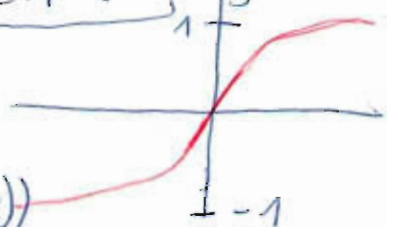
$$\alpha = 0.25$$

Coef. y func. activación bipolares

$$[x_1 = -1 \quad x_2 = 1 \quad t = 1] \quad f(x)$$

$$f(x) = \frac{1 - e^{-x}}{1 + e^{-x}}$$

$$f'(x) = \frac{1}{2} (1 + f(x))(1 - f(x))$$



$$z_{in1} = 0.4 - 0.7 - 0.2 = -0.5 \rightarrow z_1 = f(-0.5) \approx \frac{0.4}{1.4} \approx -0.29$$

$$z_{in2} = 0.6 + 0.4 + 0.3 = 1.3 \rightarrow z_2 \approx \frac{0.64}{1.3} \approx 0.57$$

$$y_{in} = -0.3 + (0.5 \cdot 0.29) + (0.1 \cdot 0.57) = -0.3 + 0.145 + 0.057 = -0.098$$

$$y = f(-0.098) \approx \frac{0.1}{1.1} \approx -0.18$$

$$\delta_y = (1 + 0.48) \cdot f'(-0.36) = 1.48 \cdot \frac{0.6 \cdot 0.9}{2} = 0.57$$

$$\delta_{in1} = 0.57 \cdot 0.5 \cdot 0.4 = 0.115 \quad \delta_1 = 0.115 \cdot f'(-0.29) = 0.13$$

$$\delta_{in2} = 0.57 \cdot 0.6 \cdot 0.3 = 0.102 \quad \delta_2 = 0.102 \cdot f'(0.57) = 0.02$$

Actualizar pesos $\Delta w_0 = 0.25 \cdot 0.57 \cdot 1 = -0.14$

$$\Delta w_{10} = 0.25 \cdot 0.57 \cdot (-0.29) = -0.04 \quad \Delta w_{02} = 0.25 \cdot 0.57 \cdot 0.57 = 0.08$$

$$\Delta v_{11} = 0.25 \cdot 0.13 \cdot (-1) = -0.03 \quad \Delta v_{12} = 0.25 \cdot 0.02 \cdot (-1) = -0.002$$

$$\Delta v_{21} = 0.25 \cdot 0.13 \cdot 1 = 0.03 \quad \Delta v_{22} = 0.25 \cdot 0.02 \cdot 1 = 0.002$$

$$v_{11} = 0.7 - 0.03 = 0.67 \quad v_{12} = -0.4 - 0.002 = -0.402$$

$$v_{21} = 0.7 + 0.03 = 0.73 \quad v_{22} = -0.4 + 0.002 = -0.398$$

$$f'(x) = \frac{1}{2} (1 + f(x))(1 - f(x))$$

$$z_{in1} = 0.4 - 0.7 - 0.2 = -0.5 \rightarrow z_1 = f(-0.5) \approx \frac{0.4}{1.4} \approx -0.29$$

$$z_{in2} = 0.6 + 0.4 + 0.3 = 1.3 \rightarrow z_2 \approx \frac{0.64}{1.3} \approx 0.57$$

$$y_{in} = -0.3 + (0.5 \cdot 0.29) + (0.1 \cdot 0.57) = -0.3 + 0.145 + 0.057 = -0.098$$

$$y = f(-0.098) \approx \frac{0.1}{1.1} \approx -0.18$$

$$\delta_y = (1 + 0.48) \cdot f'(-0.36) = 1.48 \cdot \frac{0.6 \cdot 0.9}{2} = 0.57$$

Ej. Multiple XOR

x_1	x_2	t		w_{11}	w_{12}	w_{21}	w_{22}	b_1	b_2
1	1	-1							
1	-1	1							
-1	1	1	iniz.	0.5	0.1	0.2	0.2	0.3	0.15
-1	-1	-1	11	0.5	0.1	0.2	0.2	0.3	0.15

$$v_1 = v_2 = b_3 = \frac{1}{2} \quad d = 0.5$$

1.1 $x_1 = 1 \quad x_2 = 1$

$$z_{in1} = b_1 + x_1 w_{11} + x_2 w_{21} = 0.3 + 0.5 + 0.2 = 0.55 \rightarrow 1$$

$$z_{in2} = b_2 + x_1 w_{12} + x_2 w_{22} = 0.15 + 0.1 + 0.2 = 0.45 \rightarrow 1$$

$$y_{in} = b_3 + z_1 v_1 + z_2 v_2 = 0.5 + 0.5 + 0.5 = 1.5 \rightarrow 1$$

$t = -1 \rightarrow$ actualizar pesos z_1, y, z_2

$$\Delta w_{11} = 0.5(-1 - 0.55) = -0.525 \quad \Delta w_{12} = 0.5(-1 - 0.45) = -0.475$$

$$\Delta w_{21} = 0.5(1 - 0.55) = 0.225 \quad \Delta w_{22} = 0.5(1 - 0.45) = 0.275$$

$$w_{11} = 0.5 - 0.525 = -0.025 \quad w_{12} = 0.1 - 0.475 = -0.375$$

$$\Delta b_1 = 0.5(1 - 1) = 0 \quad \Delta b_2 = 0.5(1 - 1) = 0$$

1.2 $x_1 = 1 \quad x_2 = -1$

$$z_{in1} = 0.3 + 0.5 - 0.425 = 0.375 \quad z_{in2} = 0.15 + 0.1 - 0.425 = -0.175$$

$$y_{in} = 0.5 + 0.5 - 0.5 = 0.5 \rightarrow 1 \quad t = 1 \rightarrow$$

1.3 $x_1 = -1 \quad x_2 = 1$

$$z_{in1} = 0.3 - 0.5 + 0.425 = 0.225 \quad z_{in2} = 0.15 - 0.1 + 0.425 = 0.475$$

1.4 $x_1 = 1 \quad x_2 = 1$

$$z_{in1} = b_1 + x_1 w_{11} + x_2 w_{21} = 0.3 + 0.5 + 0.2 = 0.55 \rightarrow 1$$

$$z_{in2} = b_2 + x_1 w_{12} + x_2 w_{22} = 0.15 + 0.1 + 0.2 = 0.45 \rightarrow 1$$

$$y_{in} = b_3 + z_1 v_1 + z_2 v_2 = 0.5 + 0.5 + 0.5 = 1.5 \rightarrow 1$$

$t = -1 \rightarrow$ actualizar pesos z_1, y, z_2

$$\Delta w_{11} = 0.5(-1 - 0.55) = -0.525 \quad \Delta w_{12} = 0.5(-1 - 0.45) = -0.475$$

$$\Delta w_{21} = 0.5(1 - 0.55) = 0.225 \quad \Delta w_{22} = 0.5(1 - 0.45) = 0.275$$

Autoasoc. con func. umbral

FASE 1: "entrenar" la red para q. almacene el vector $V = (1, 1, 1, -1)$

$$\begin{pmatrix} 1 \\ 1 \\ 1 \\ -1 \end{pmatrix} (1 \ 1 \ 1 \ -1) = \begin{pmatrix} 1 & 1 & 1 & -1 \\ 1 & 1 & 1 & -1 \\ 1 & 1 & 1 & -1 \\ -1 & -1 & -1 & 1 \end{pmatrix} \xrightarrow[\text{autoconexiones}]{\text{eliminar}} W = \begin{pmatrix} 0 & 1 & 1 & -1 \\ 1 & 0 & 1 & -1 \\ 1 & 1 & 0 & -1 \\ -1 & -1 & -1 & 0 \end{pmatrix}$$

FASE 2: Probar la red con los vectores $(1 \ 0 \ 0 \ 0)$ y $(0 \ 0 \ 0 \ -1)$

$\theta = 0$

$$x_0 = (1 \ 0 \ 0 \ 0) \quad x_1 = x_0 \cdot W = (0 \ 1 \ 1 \ -1)$$

$$x_2 = x_1 \cdot W = (3 \ 2 \ 2 \ -2) \rightarrow (1 \ 1 \ 1 \ -1) = V \rightarrow \text{stop}$$

$$x_0 = (0 \ 0 \ 0 \ -1) \quad x_1 = x_0 \cdot W = (1 \ 1 \ 1 \ 0)$$

$$x_2 = x_1 \cdot W = (2 \ 2 \ 2 \ -3) = (1 \ 1 \ 1 \ -1)$$

$$x_2 = V \rightarrow \text{stop}$$

La red ha reconocido a los 2 vectores, cada uno de los cuales tenía 3 ~~errores~~ elementos ~~erróneos~~ desconocidos.

Probar con un vector con 2 errores: $V_3 = (-1 \ -1 \ 1 \ -1)$

$$x_0 = (-1 \ -1 \ 1 \ -1) \quad x_1 = x_0 \cdot W = (1 \ 1 \ -1 \ 1)$$

$$x_2 = x_1 \cdot W = (-1 \ -1 \ 1 \ -1) = x_0 \quad \text{Stop}$$

FASE 2: Probar la red con los vectores $(1 \ 0 \ 0 \ 0)$ y $(0 \ 0 \ 0 \ -1)$

$\theta = 0$

$$x_0 = (1 \ 0 \ 0 \ 0) \quad x_1 = x_0 \cdot W = (0 \ 1 \ 1 \ -1)$$

$$x_2 = x_1 \cdot W = (3 \ 2 \ 2 \ -2) \rightarrow (1 \ 1 \ 1 \ -1) = V \rightarrow \text{stop}$$

$$x_0 = (0 \ 0 \ 0 \ -1) \quad x_1 = x_0 \cdot W = (1 \ 1 \ 1 \ 0)$$

led autoasoc. ~~Real~~ ~~autoasoc~~

Obtener la matriz de pesos para almacenar los sigtes. vectores (ortogonales)

$$(1 \ 1 \ 1 \ 1)$$

$$(1 \ 1 \ -1 \ -1)$$

$$(1 \ -1 \ 1 \ -1)$$

$$(1 \ -1 \ -1 \ 1)$$

$$\begin{matrix} a_{11} \\ a_{21} \\ a_{31} \\ a_{41} \end{matrix} \begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix} \quad \begin{matrix} b_{11} & b_{12} & b_{13} & b_{14} \\ (1 & 1 & 1 & 1) \end{matrix} = \left(\begin{matrix} 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \end{matrix} \right) \left\{ \begin{pmatrix} 1 \\ 1 \\ -1 \\ -1 \end{pmatrix} (1 \ 1 \ -1 \ -1) = \begin{pmatrix} 1 & 1 & -1 & -1 \\ 1 & 1 & -1 & -1 \\ -1 & -1 & 1 & 1 \\ -1 & -1 & 1 & 1 \end{pmatrix} \right.$$

$$4 \times 1 \cdot 1 \times 4 \rightarrow 4 \times 4$$

$$K, l \rightarrow a_{kl} b_{1l} + a_{k2} b_{2l} + \dots + a_{kn} b_{nl}$$

$$1, 1 \rightarrow a_{11} b_{11} + a_{12} b_{21} + a_{13} b_{31}$$

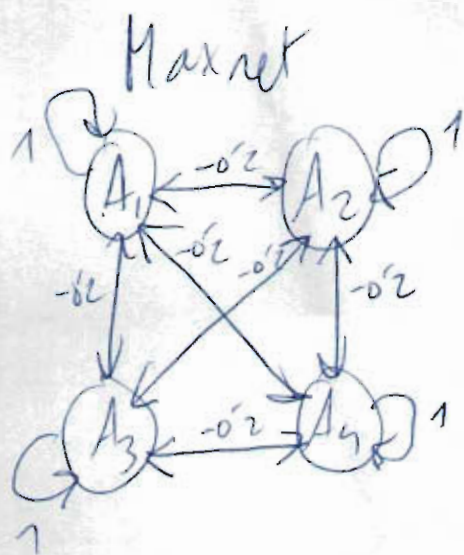
$$\begin{pmatrix} 1 \\ -1 \\ 1 \\ -1 \end{pmatrix} (1 \ -1 \ 1 \ -1) = \begin{pmatrix} 1 & -1 & 1 & -1 \\ -1 & 1 & -1 & 1 \\ 1 & -1 & 1 & -1 \\ -1 & 1 & -1 & 1 \end{pmatrix} \left\{ \begin{pmatrix} 1 \\ -1 \\ -1 \\ 1 \end{pmatrix} (1 \ 1 \ -1 \ 1) = \begin{pmatrix} 1 & -1 & -1 & 1 \\ -1 & 1 & 1 & -1 \\ -1 & 1 & 1 & -1 \\ 1 & -1 & -1 & 1 \end{pmatrix} \right.$$

$$W_1 + W_2 = \begin{pmatrix} 2 & 2 & 0 & 0 \\ 2 & 0 & 0 & 0 \\ 0 & 0 & 2 & 2 \\ (1 & -1 & -1 & 1) \end{pmatrix} \quad \text{Comprobamos si cumplen los vectores 1 y 2:}$$

$$\begin{matrix} a_{11} \\ a_{21} \\ a_{31} \\ a_{41} \end{matrix} \begin{pmatrix} 1 \\ 1 \\ 1 \\ 1 \end{pmatrix} \quad \begin{matrix} b_{11} & b_{12} & b_{13} & b_{14} \\ (1 & 1 & 1 & 1) \end{matrix} = \left(\begin{matrix} 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \end{matrix} \right) \left\{ \begin{pmatrix} 1 \\ 1 \\ -1 \\ -1 \end{pmatrix} (1 \ 1 \ -1 \ -1) = \begin{pmatrix} 1 & 1 & -1 & -1 \\ 1 & 1 & -1 & -1 \\ -1 & -1 & 1 & 1 \\ -1 & -1 & 1 & 1 \end{pmatrix} \right.$$

$$4 \times 1 \cdot 1 \times 4 \rightarrow 4 \times 4$$

$$K, l \rightarrow a_{kl} b_{1l} + a_{k2} b_{2l} + \dots + a_{kn} b_{nl}$$



$$0 \quad x = (0.2 \quad 0.4 \quad 0.6 \quad 0.8)$$

$$1 \quad x_1(1) = (x_1(0) - 0.2(0.4 + 0.6 + 0.8)) = -0.07 \rightarrow 0$$

$$x_2(1) = (0.4 - 0.2(0.2 + 0.6 + 0.8)) = 0.08 \rightarrow 0.08$$

$$x_3(1) = (0.6 - 0.2(0.2 + 0.4 + 0.8)) = 0.32 \rightarrow 0.32$$

$$x_4(1) = (0.8 - 0.2(0.2 + 0.4 + 0.6)) = 0.56 \rightarrow 0.56$$

$$x = (0 \quad 0.08 \quad 0.32 \quad 0.56)$$

$$2 \quad x_1(2) = 0 - K = -K \rightarrow 0$$

$$x_2(2) = (0.08 - 0.2(0.32 + 0.56)) = -0.168 \rightarrow 0$$

$$x_3(2) = (0.32 - 0.2(0 + 0.08 + 0.56)) = 0.192$$

$$x_4(2) = (0.56 - 0.2(0 + 0.08 + 0.32)) = 0.48$$

$$x = (0 \quad 0 \quad 0.192 \quad 0.48)$$

$$3 \quad x_3(3) = (0.192 - 0.2(0.48)) = 0.096$$

$$x_4(3) = (0.48 - 0.2 \cdot 0.192) = 0.442$$

$$x = (0 \quad 0 \quad 0.096 \quad 0.442)$$

$$4 \quad x = (0 \quad 0 \quad 0.008 \quad 0.422)$$

$$5 \quad x = (0 \quad 0 \quad 0 \quad 0.421)$$

$$x = (0 \quad 0.08 \quad 0.32 \quad 0.56)$$

$$2 \quad x_1(2) = 0 - K = -K \rightarrow 0$$

$$x_2(2) = (0.08 - 0.2(0.32 + 0.56)) = -0.168 \rightarrow 0$$

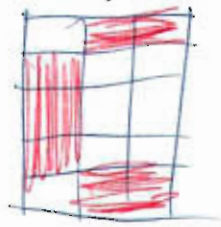
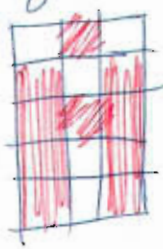
$$x_3(2) = (0.32 - 0.2(0 + 0.08 + 0.56)) = 0.192$$

$$x_4(2) = (0.56 - 0.2(0 + 0.08 + 0.32)) = 0.48$$

$$x = (0 \quad 0 \quad 0.192 \quad 0.48)$$

BAM

Se quiere asociar los letreros A y C con los códigos $(-1, 1)$ y $(1, 1)$, respectivamente



Los letreros se codifican bipolares. P.e., la ~~matriz~~ para el vector para A ser:

$(-1, 1, -1, 1, -1, 1, 1, 1, 1, 1, -1, 1, 1, -1, 1)$

FASE 1: inicializar pesos red:

$$\begin{aligned}
 [A] \begin{pmatrix} -1 \\ 1 \\ -1 \\ 1 \\ -1 \\ 1 \\ 1 \\ 1 \\ 1 \\ -1 \\ 1 \\ 1 \\ -1 \\ 1 \\ 1 \end{pmatrix} & \begin{pmatrix} 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \end{pmatrix} & [C] \begin{pmatrix} -1 \\ 1 \\ 1 \\ 1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \end{pmatrix} & \begin{pmatrix} -1 & -1 \\ 1 & 1 \\ 1 & 1 \\ 1 & 1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \end{pmatrix} & W = \begin{pmatrix} 0 & -2 \\ 0 & 2 \\ 2 & 0 \\ 0 & 2 \\ 0 & -2 \\ -2 & 0 \\ 0 & 2 \\ -2 & 0 \\ -2 & 0 \\ 0 & 2 \\ 0 & -2 \\ -2 & 0 \\ -2 & 0 \\ 2 & 0 \\ 0 & 2 \end{pmatrix} \\
 & (-1 \ 1) = & (1 \ 1) = & \pm W_1 + W_2 =
 \end{aligned}$$

FASE 2: recuperación info.

$$\text{En } (1, 1) \rightarrow x_{in} = y W^T = [-2 \ 2 \ -2 \ 2 \ -2 \ 2 \ 2 \ 2 \ 2 \ 2 \ -2 \ 2 \ 2 \ -2 \ 2]$$

Vemos q a partir del código de A se recupera la propia A

$(-1, 1, -1, 1, -1, 1, 1, 1, 1, 1, -1, 1, 1, -1, 1)$

FASE 1: inicializar pesos red:

$$\begin{aligned}
 [A] \begin{pmatrix} -1 \\ 1 \\ -1 \\ 1 \\ -1 \\ 1 \\ 1 \\ 1 \\ 1 \\ -1 \\ 1 \\ 1 \\ -1 \\ 1 \\ 1 \end{pmatrix} & \begin{pmatrix} 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \\ -1 & 1 \\ 1 & -1 \end{pmatrix} & [C] \begin{pmatrix} -1 \\ 1 \\ 1 \\ 1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \\ -1 \end{pmatrix} & \begin{pmatrix} -1 & -1 \\ 1 & 1 \\ 1 & 1 \\ 1 & 1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \\ -1 & -1 \end{pmatrix} & W = \begin{pmatrix} 0 & -2 \\ 0 & 2 \\ 2 & 0 \\ 0 & 2 \\ 0 & -2 \\ -2 & 0 \\ 0 & 2 \\ -2 & 0 \\ -2 & 0 \\ 0 & 2 \\ 0 & -2 \\ -2 & 0 \\ -2 & 0 \\ 2 & 0 \\ 0 & 2 \end{pmatrix} \\
 & (-1 \ 1) = & (1 \ 1) = & \pm W_1 + W_2 =
 \end{aligned}$$

Sombrero mexicano

$$f(x) = \begin{cases} 0 & x < 0 \\ x & 0 \leq x \leq 2 \\ 2 & 2 < x \end{cases}$$

$$R_1 = 1, C_1 = 0.6$$

$$R_2 = 2, C_2 = -0.4$$

① ② ③ ④ ⑤ ⑥ ⑦

$$t=0 \quad X = (0, 0.5, 0.8, 1, 0.8, 0.5, 0)$$

$$t=1 \quad \begin{aligned} x_{in1}(1) &= 0.6(0 + 0.5) - 0.4 \cdot 0.8 = -0.02 \rightarrow 0 \\ x_{in2}(1) &= 0.6(0 + 0.5 + 0.8) - 0.4 \cdot 1 = 0.38 \rightarrow 0.38 \\ x_{in3}(1) &= 0.6(0.5 + 0.8 + 1) - 0.4(0 + 0.8) = 1.06 \rightarrow 1.06 \\ x_{in4}(1) &= 0.6(0.8 + 1 + 0.8) - 0.4(0.5 + 0.5) = 1.16 \rightarrow 1.16 \\ x_{in5}(1) &= 0.6(1 + 0.8 + 0.5) - 0.4(0.8 + 0) = 1.06 \rightarrow 1.06 \\ x_{in6}(1) &= 0.6(0.8 + 0.5 + 0) - 0.4 \cdot 1 = 0.38 \rightarrow 0.38 \\ x_{in7}(1) &= 0.6(0.5 + 0) - 0.4 \cdot 0.8 = -0.02 \rightarrow 0 \end{aligned}$$

$$X = (0 \ 0.38 \ 1.06 \ 1.16 \ 1.06 \ 0.38 \ 0)$$

$$t=2 \quad \begin{aligned} x_{in1}(2) &= 0.6(0 + 0.38) - 0.4 \cdot 1.06 = -0.196 \rightarrow 0 \\ x_{in2}(2) &= 0.6(0 + 0.38 + 1.06) - 0.4(1.16) = 0.39 \rightarrow 0.39 \\ x_{in3}(2) &= 0.6(0.38 + 1.06 + 1.16) - 0.4(0 + 1.06) = 1.14 \rightarrow 1.14 \\ x_{in4}(2) &= 0.6(1.06 + 1.16 + 1.06) - 0.4(0.38 + 0.38) = 1.66 \rightarrow 1.66 \\ x_{in5}(2) &= 1.136 \rightarrow 1.14 \\ x_{in6}(2) &= 0.39 \rightarrow 0.39 \\ x_{in7}(2) &= -0.196 \rightarrow 0 \end{aligned}$$

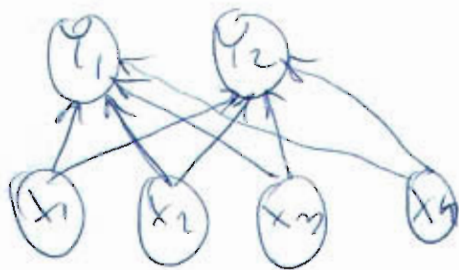
$$t=1 \quad \begin{aligned} x_{in1}(1) &= 0.6(0 + 0.38) - 0.4 \cdot 1.06 = -0.02 \rightarrow 0 \\ x_{in2}(1) &= 0.6(0 + 0.5 + 0.8) - 0.4 \cdot 1 = 0.38 \rightarrow 0.38 \\ x_{in3}(1) &= 0.6(0.5 + 0.8 + 1) - 0.4(0 + 0.8) = 1.06 \rightarrow 1.06 \\ x_{in4}(1) &= 0.6(0.8 + 1 + 0.8) - 0.4(0.5 + 0.5) = 1.16 \rightarrow 1.16 \\ x_{in5}(1) &= 0.6(1 + 0.8 + 0.5) - 0.4(0.8 + 0) = 1.06 \rightarrow 1.06 \\ x_{in6}(1) &= 0.6(0.8 + 0.5 + 0) - 0.4 \cdot 1 = 0.38 \rightarrow 0.38 \\ x_{in7}(1) &= 0.6(0.5 + 0) - 0.4 \cdot 0.8 = -0.02 \rightarrow 0 \end{aligned}$$

$$X = (0 \ 0.38 \ 1.06 \ 1.16 \ 1.06 \ 0.38 \ 0)$$

Hamming

Vectores ejemplares: $e(1) = (1, -1, -1, -1)$
 $e(2) = (-1, -1, -1, 1)$

Almacenar los vectores en la red



$$w_{ij} = \frac{e_i(l_j)}{2}$$

$$W = \begin{pmatrix} 0's & -0's \\ -0's & -0's \\ -0's & -0's \\ -0's & 0's \end{pmatrix}$$

Bias $\rightarrow b_1 = b_2 = \frac{4}{2} = 2$

Vectores de entrada

$$x = (1, 1, -1, -1) \rightarrow y_{in1} = b_1 + \sum_1^4 x_i w_{i1} = 2 + (0's - 0's + 0's + 0's) = 3$$

$$y_{in2} = 2 + (-0's - 0's + 0's - 0's) = 1$$

~~Aplicar Maxnet~~ $y = y_{in} = (3, 1)$

Aplicar Maxnet $\rightarrow y_1 \rightarrow x$ se parece a y_1

$$x = (-1, -1, 1, 1) \rightarrow y_{in1} = 2 + (-0's + 0's - 0's - 0's) = 1$$

$$y_{in2} = 2 + (0's + 0's + 0's + 0's) = 3$$

$$y = y_{in} = (1, 3)$$

Maxnet $\rightarrow y_2 \rightarrow x$ se parece más a y_2



$$W = \begin{pmatrix} -0's & -0's \\ -0's & 0's \end{pmatrix}$$

Bias $\rightarrow b_1 = b_2 = \frac{4}{2} = 2$

Vectores de entrada

$$x = (1, 1, -1, -1) \rightarrow y_{in1} = b_1 + \sum_1^4 x_i w_{i1} = 2 + (0's - 0's + 0's + 0's) = 3$$

$$y_{in2} = 2 + (-0's - 0's + 0's - 0's) = 1$$

Hopfield discrete

Red q. almacena el vector $(1, 1, 1, 0)$. $\theta_i = 0$
 Aprendizaje del vector (no supervisado). Se pasa el vector a bipolar $(1, 1, 1, -1)$

$$\begin{pmatrix} 1 \\ 1 \\ 1 \\ -1 \end{pmatrix} (1, 1, 1, -1) = \begin{pmatrix} x_0 & 1 & 1 & -1 \\ 1 & x_0 & 1 & -1 \\ 1 & 1 & x_0 & -1 \\ -1 & -1 & -1 & x_0 \end{pmatrix} = W \quad \leftarrow \text{diagonal a 0}$$

(si hubiera + patrones, habría q. sumar los respectivos patrones)

Vector entrada $\rightarrow x = (0, 0, 1, 0)$ 2 valores de ~~errores~~ de ~~sanados~~

(0) Elección aleatoria de unid. $\rightarrow Y_1$

$$y_{in1} = 0 + (0 + 0 + 1 + 0) = 1 \rightarrow 1$$

$$y = (1, 0, 1, 0)$$

(1) Elección $\rightarrow Y_4$

$$y_{in4} = 0 + (-1 + 0 - 1 + 0) = -2 \rightarrow 0$$

$$y = (1, 0, 1, 0)$$

(2) Elección $\rightarrow Y_3$

$$y_{in3} = 1 + (1 + 0 + 0 + 0) = 2 \rightarrow 1$$

$$y = (1, 0, 1, 0)$$

$$\begin{pmatrix} 1 \\ 1 \\ 1 \\ -1 \end{pmatrix} \begin{pmatrix} 1 \\ 1 \\ 1 \\ -1 \end{pmatrix} = \begin{pmatrix} x_0 & 1 & 1 & -1 \\ 1 & x_0 & 1 & -1 \\ 1 & 1 & x_0 & -1 \\ -1 & -1 & -1 & x_0 \end{pmatrix} \quad \leftarrow \text{diagonal a 0}$$

(si hubiera + patrones, habría q. sumar los respectivos patrones)

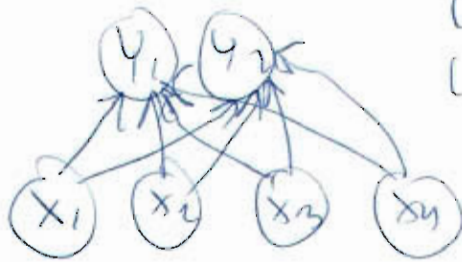
Vector entrada $\rightarrow x = (0, 0, 1, 0)$ 2 valores de ~~errores~~ de ~~sanados~~

(0) Elección aleatoria de unid. $\rightarrow Y_1$

$$y_{in1} = 0 + (0 + 0 + 1 + 0) = 1 \rightarrow 1$$

SOM

Vectores de entrada: $(1, 1, 0, 0)$
 $(0, 0, 0, 1)$
 $(1, 0, 0, 0)$
 $(0, 0, 1, 1)$



$$\alpha(0) = 0.6$$

$$\alpha(t+1) = 0.5 \cdot \alpha(t)$$

Radio 0

Sin vecindad

$N^{\circ} \text{máx. clusters} = 2$

1) Inicializ. matriz W (aleatorio) $W = \begin{pmatrix} 0.2 & 0.8 \\ 0.6 & 0.4 \\ 0.5 & 0.7 \\ 0.9 & 0.3 \end{pmatrix}$

$x = (1, 1, 0, 0)$ Buscar neurona + próxima

$$\alpha = 0.6$$

$$D(1) = (0.2 - 1)^2 + (0.6 - 1)^2 + 0.5^2 + 0.9^2 = 1.86$$

$$\rightarrow D(2) = (0.8 - 1)^2 + (0.4 - 1)^2 + 0.7^2 + 0.3^2 = 0.98$$

Actualizar pesos de y_2

$$\Delta w_{i2} = \alpha(x_i - w_{i2}(\text{old}))$$

$$\Delta w_{12} = 0.6(1 - 0.8) = 0.12$$

$$\Delta w_{22} = 0.6(1 - 0.4) = 0.36$$

$$\Delta w_{32} = 0.6(0 - 0.7) = -0.42$$

$$\Delta w_{42} = 0.6(0 - 0.3) = -0.18$$

$$W^T = \begin{pmatrix} 0.2 & 0.6 & 0.5 & 0.9 \\ 0.42 & 0.4 & 0.28 & 0.12 \end{pmatrix}$$

$$x = (0, 0, 0, 1) \rightarrow D(1) = (0.2 - 0)^2 + 0.6^2 + 0.5^2 + (0.9 - 1)^2 = 0.66$$

$$\alpha = 0.6$$

$$D(2) = 0.92^2 + 0.76^2 + 0.28^2 + (0.12 - 1)^2 = 2.27$$

$$\Delta w_{11} = 0.6(0 - 0.2) = -0.12$$

$$\Delta w_{21} = 0.6(0 - 0.6) = -0.36$$

$$\Delta w_{31} = 0.6(0 - 0.5) = -0.30$$

$$\Delta w_{41} = 0.6(1 - 0.9) = 0.06$$

$$W^T = \begin{pmatrix} 0.08 & 0.24 & 0.2 & 0.96 \\ 0.28 & 0.24 & 0.22 & 0.12 \end{pmatrix}$$

1) Inicializ. matriz W (aleatorio) $W = \begin{pmatrix} 0.2 & 0.8 \\ 0.6 & 0.4 \\ 0.5 & 0.7 \\ 0.9 & 0.3 \end{pmatrix}$

$x = (1, 1, 0, 0)$ Buscar neurona + próxima

$$\alpha = 0.6$$

$$D(1) = (0.2 - 1)^2 + (0.6 - 1)^2 + 0.5^2 + 0.9^2 = 1.86$$

$$\rightarrow D(2) = (0.8 - 1)^2 + (0.4 - 1)^2 + 0.7^2 + 0.3^2 = 0.98$$

Actualizar pesos de y_2

$$\Delta w_{i2} = \alpha(x_i - w_{i2}(\text{old}))$$

$$X = (0, 0, 1, 1) \rightarrow D(1) = (0'08)^2 + (0'24)^2 + (0'2-1)^2 + (0'98-1)^2 = 0'0064 + \dots = 0'70$$

$$D(2) = 0'97^2 + 0'3^2 + (0'12-1)^2 + (0'05-1)^2 = 0'94 + 0'09 + 0'77 + 0'9 = 2'70$$

$$\Delta w_{11} = 0'6(0 - 0'08) = -0'048 \quad \Delta w_{21} = 0'6(0 - 0'24) = -0'144$$

$$\Delta w_{31} = 0'6(1 - 0'2) = 0'48 \quad \Delta w_{41} = 0'6(1 - 0'98) = 0'012$$

$$W^+ = \begin{pmatrix} 0'32 & 0'09 & 0'68 & 0'98 \\ 0'97 & 0'3 & 0'12 & 0'05 \end{pmatrix}$$

Después de esta época se actualiza $\alpha = 0'3$

Después de la 2ª época, la matriz es

$$W(2) = \begin{pmatrix} 0'016 & 0'48 \\ 0'047 & 0'36 \\ 0'630 & 0'05 \\ 0'991 & 0'024 \end{pmatrix}$$

$$W(100) = \begin{pmatrix} 0 & 1 \\ 0 & 0'49 \\ 0'51 & 0 \\ 1 & 0 \end{pmatrix} \begin{matrix} \rightarrow x_{1,3} \\ \rightarrow x_{2,4} \end{matrix}$$

LVQ

Vectores

Clase

(1, 1, 0, 0)

1

(0, 0, 0, 1)

2

(0, 0, 1, 1)

2

(1, 0, 0, 0)

1

$$W^+ = \begin{pmatrix} 0'32 & 0'09 & 0'68 & 0'98 \\ 0'97 & 0'3 & 0'12 & 0'05 \end{pmatrix}$$

$$W = \begin{pmatrix} 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

$$X = (0, 0, 1, 0)$$

$$D(1) = 1^2 + 1^2 + 1^2 + 1^2 = 4$$

$$D(2) = \|0 - 0, 1 - 0\| = 1 \leftarrow$$

Clase esperada = clase obtenida $\rightarrow$ Sumar (reducir)

Después de esta época se actualiza $\alpha = 0'3$

Después de la 2ª época, la matriz es

$$W(2) = \begin{pmatrix} 0'016 & 0'48 \\ 0'047 & 0'36 \\ 0'630 & 0'05 \\ 0'991 & 0'024 \end{pmatrix}$$

$$W(100) = \begin{pmatrix} 0 & 1 \\ 0 & 0'49 \\ 0'51 & 0 \\ 1 & 0 \end{pmatrix} \begin{matrix} \rightarrow x_{1,3} \\ \rightarrow x_{2,4} \end{matrix}$$

clase esperada = clase obtenida

$$w_1(\text{new}) = (1 \ 1 \ 0 \ 0) + 0'1 (0 \ -1 \ 0 \ 0) = (1 \ 0'9 \ 0 \ 0)$$

$$W = \begin{pmatrix} 1 & 0'9 & 0 & 0 \\ 0 & 0 & 0'1 & 1 \end{pmatrix}$$

$$X = (0 \ 1 \ 1 \ 0)$$

$$P(1) = \| 1 \ -0'1 \ -1 \ 0 \| \approx \sqrt{2} \leftarrow$$

$$P(2) = \| 0 \ -1 \ -0'9 \ 1 \| \approx \sqrt{3}$$

clase esperada $\neq$ clase obtenida

$$w_1(\text{new}) = (1 \ 0'9 \ 0 \ 0) - 0'1 (-1 \ 0'1 \ 1 \ 0) = (1'1 \ 0'89 \ -0'1 \ 0)$$

$$W = \begin{pmatrix} 1'1 & 0'89 & -0'1 & 0 \\ 0 & 0 & 0'1 & 1 \end{pmatrix}$$

thus converges to 2nd epoch.

$$P(2) = \| 0 \ -1 \ -0'9 \ 1 \| \approx \sqrt{3}$$

clase esperada $\neq$ clase obtenida

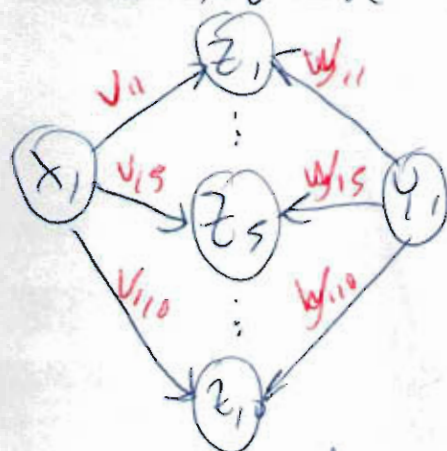
$$w_1(\text{new}) = (1 \ 0'9 \ 0 \ 0) - 0'1 (-1 \ 0'1 \ 1 \ 0) = (1'1 \ 0'89 \ -0'1 \ 0)$$

$$W = \begin{pmatrix} 1'1 & 0'89 & -0'1 & 0 \\ 0 & 0 & 0'1 & 1 \end{pmatrix}$$

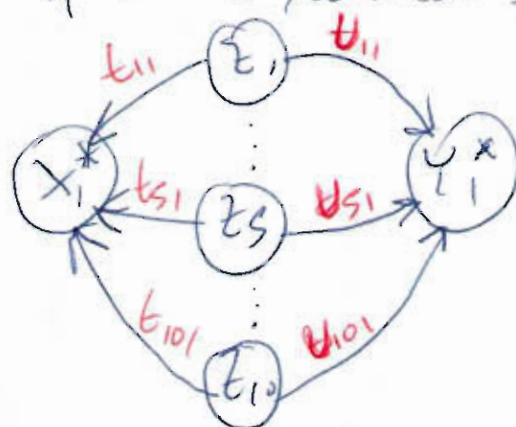
thus converges to 2nd epoch.

Contrapropagación

Se desea un aproximado de la func. $f(x) = \frac{1}{x}$
 Se tienen 1000 pto. en $[0.01, 10]$ para entrenar la red. Se estima q. con 10 neuronas en la capa



Red en la fase 1



Red en la fase 2

de cluster es suficiente. Para la capa de entrada y salida sólo se necesita 1 ~~cluster~~ neurona.

$$X=10 \quad y=f(x)=0.1$$

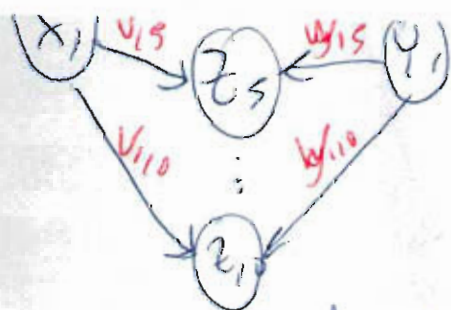
FASE 1

1) Inicializ. pesos a valores aleatorios dentro de $[0.01, 10]$

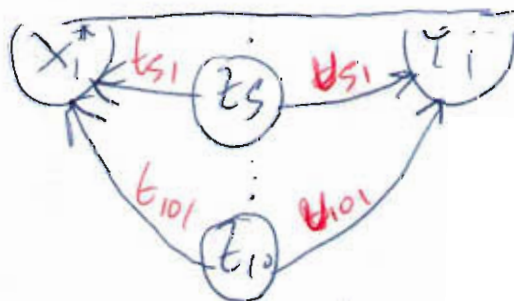
$$v = (0.2 \quad 3.1 \quad 5.2 \quad 6.4 \quad 2.2 \quad 8.1 \quad 1.5 \quad 9.1 \quad 4.4 \quad 9.9) \\ w = (9.9 \quad 4.4 \quad 9.1 \quad 1.5 \quad 8.1 \quad 2.2 \quad 6.4 \quad 5.2 \quad 3.1 \quad 0.2) \\ t = (0.3 \quad \dots \quad 8) \quad u = (7 \quad \dots \quad 2)$$

2) Para cada patrón de entrada ajustar los pesos

$$x=10 \quad y=0.1 \quad D_1 = \sqrt{(10-0.2)^2 + (0.1-9.9)^2} \approx 15$$



Red en la fase 1



Red en la fase 2

de cluster es suficiente. Para la capa de entrada y salida sólo se necesita 1 ~~cluster~~ neurona.

Unid. del cluster	z_1	z_2	z_3	z_4	z_5	z_6	z_7	z_8	z_9	z_{10}
v	0'11	0'14	0'2	0'3	0'6	1'6	3'3	5	7	9
w	9	7	5	3'3	1'6	0'6	0'3	0'2	0'14	0'11

FASE 2 Se pasan de nuevo los vectores de entrenamiento

$$x=1 \quad y=1 \quad p(1) = \sqrt{(1-0'11)^2 + (1-9)^2} \approx 8$$

$$\rightarrow p(5) = \sqrt{(1-0'6)^2 + (1-1'6)^2} \approx 0'2$$

$$p(6) = \approx 0'2$$

$$p(10) = \sqrt{(1-9)^2 + (1-0'11)^2} \approx 8$$

la unid. ganadora es z_5 . Se actualizan los pesos:

$$\Delta v_{15} = 0'005 \cdot (1-0'6) = 0'002 \quad v_{15} = 0'602$$

$$\Delta w_{15} = 0'005(1-1'6) = 0'003 \quad w_{15} = 1'603$$

$$\Delta t_{51} = 0'5(1-2) = -0'5 \quad t_{51} = 1'5$$

$$\Delta u_{51} = 0'5(1-0'5) = 0'25 \quad u_{51} = 0'25$$

Actualizar tasas aprendizaje: $\alpha = 0'004 = \beta$; $a = 0'49 = b$
 Tras pasar los 1000 patrones, los pesos quedan así:

v	0'11	0'14	0'2	0'3	0'6	1'6	3'3	5	7	9
w	9	7	5	3'3	1'6	0'6	0'3	0'2	0'14	0'11
t	0'11	0'14	0'2	0'3	0'6	1'6	3'3	5	7	9
u	9	7	5	3	1'6	0'6	0'3	0'2	0'14	0'11

$$\rightarrow p(5) = \sqrt{(1-0'6)^2 + (1-1'6)^2} \approx 0'2$$

$$p(6) = \approx 0'2$$

$$p(10) = \sqrt{(1-9)^2 + (1-0'11)^2} \approx 8$$

la unid. ganadora es z_5 . Se actualizan los pesos:

$$\Delta v_{15} = 0'005 \cdot (1-0'6) = 0'002 \quad v_{15} = 0'602$$

$$\Delta w_{15} = 0'005(1-1'6) = 0'003 \quad w_{15} = 1'603$$

$$\Delta t_{51} = 0'5(1-2) = -0'5 \quad t_{51} = 1'5$$

1) Buscar la uid. del cluster + próxima

$$D_1 = \sqrt{(0'12 - 0'11)^2 + (0 - 9)^2} \simeq 9$$

⋮

$$\rightarrow D_6 = \sqrt{(0'12 - 1'6)^2 + (0 - 0'6)^2} \simeq 1'6$$

⋮

$$D_{10} = \sqrt{(0'12 - 9)^2 + (0 - 0'11)^2} \simeq 8'9$$

2) Calcular la aprox.:

$$x^* = t_6 = 1'6 \quad y^* = 0'6$$

La aprox. es burda, pq. hemos usado un valor desconocido de "y" para la búsqueda. Si buscamos en el cluster sólo con x,

$$\rightarrow D_1 = \sqrt{(0'12 - 0'11)^2} = 0'01$$

$$D_6 = \sqrt{(0'12 - 1'6)^2} = 1'48$$

⋮

$$D_{10} = \sqrt{(0'12 - 9)^2} = 8'88$$

Calculando la aprox. con z_1 ,

$$\boxed{\begin{array}{l} x^* = t_1 = 0'11 \\ y^* = u_1 = 9 \end{array}}$$

2) Calcular la aprox.:

$$x^* = t_6 = 1'6 \quad y^* = 0'6$$

La aprox. es burda, pq. hemos usado un valor desconocido de "y" para la búsqueda. Si buscamos en el cluster sólo con x,

$$\rightarrow D_1 = \sqrt{(0'12 - 0'11)^2} = 0'01$$

$$D_6 = \sqrt{(0'12 - 1'6)^2} = 1'48$$