



Faculteit der Exacte Wetenschappen

(Duidelijk en met blokletters invullen)

1 2 3 4 5 6 7 8 9 10
155 | 10 | 10 | 5 | 10 | 10 | 10 | 10 | 10

Naam en voorletters: .

Cijfer:

Vak: Neurale netwerken

Studentnummer: .

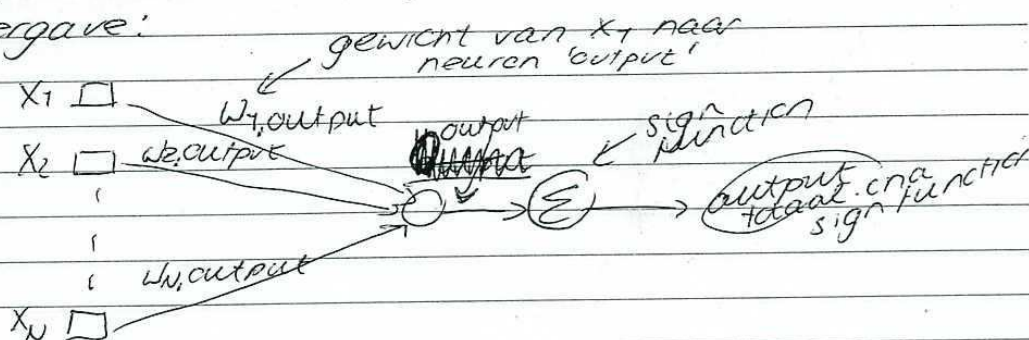
Datum: 26-05-05

Jaar van 1e inschrijving: 99

Studierichting: BWL

vrije Universiteit amsterdam

1. Perceptron architectuur: architectuur van een perceptron bestaat uit één enkele laag (single layer). Dit houdt in dat de input neuronen rechtstreeks verbonden zijn met de output neuron. Een perceptron heeft maar één output neuron. Grafische weergave:



- Neuron model (information processing unit): in ~~een~~ ^{de} ~~betere~~ ^{output} neuron wordt de sign functie gebruikt. Deze is als volgt gedefinieerd:

$$f(v_j) = \begin{cases} 1 & \text{als } v_j \geq 0 \\ -1 & \text{als } v_j < 0 \end{cases}$$

$$v_j = \sum w_{ij} x_i$$

- learning algoritme:

Voor de perceptron wordt het fixed increment learning algoritme gebruikt:

$n=1$

initialize $w(n)$ random.

while (there are misclassified examples)

select a misclassified example $(x(n), d(n))$

$$w(n+1) = w(n) + \eta x(n) d(n)$$

$n=n+1$

end- while.

Hierbij is n het de oewichtvector. n het aantal

aanpassingen dat je doet op de gewichtsvector w de learning rate en $x(n)$ het training example dat je voor de n^e iteratie gebruikt en $d(n)$ de ~~class~~ (correcte) classificatie van dit n^e training example.

• ~~Perceptron kan niet worden toegepast~~

ervolg vraag 1.

speseparable classes

5 • Perceptron kan niet worden toegepast op non-linearly training examples omdat perceptron alleen kan classificeren d.m.v. één hyperplane (rechte lijn / vlak) als examples niet linear separabel zijn, zijn er dus bijvoorbeeld meerdere lijnen / vlakken nodig en dus kan de perceptron dit probleem niet aan.

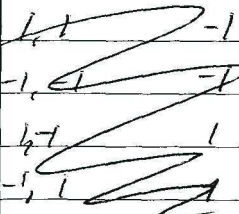
* Als examples niet linear separabel zijn, zijn er meerdere lijnen / vlakken nodig om de examples te scheiden en de perceptron kan het slechts met één hyperplane.

2. ~~Supervised learning~~: je input bestaat uit ~~clabeled input~~

2 • Supervised learning: je data bestaat uit
 10 clabeled inputs, desired outputs. Dit houdt in dat je bij het trainen van je netwerk dus niet alleen de beschikking hebt over labeled inputs (dus de input voor je netwerk) maar je weet ook wat het netwerk bij een bepaalde input als output (=desired output) moet geven. Deze desired outputs gebruik je ook bij het trainen van je netwerk. Voorbeelden van supervised learning zijn perceptron, adaline, feed forward neural networks en RBF ~~radial basis functions~~.

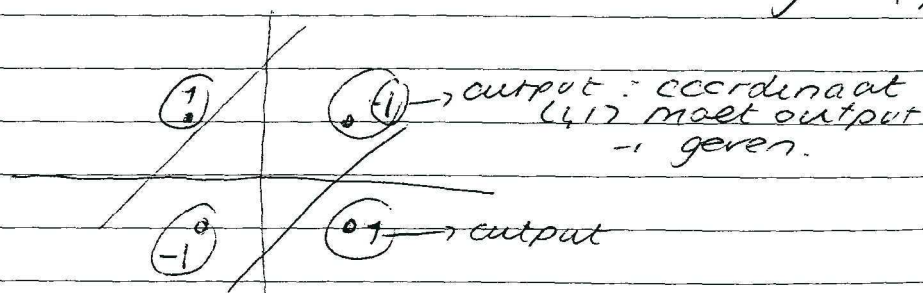
Unsupervised learning: Bij unsupervised learning heb je alleen beschikking over inputs en weet je dus niet welke output er bij een dergelijke (set) input hoort. Je gaat het netwerk trainen ~~op basis van~~ met behulp van ^{alleen} je inputs. (Je kunt bijvoorbeeld structuur in je inputs ontdekken). Voorbeelden van unsupervised learning zijn SOM en Hopfield.

• a FFNN that solves the XOR problem.

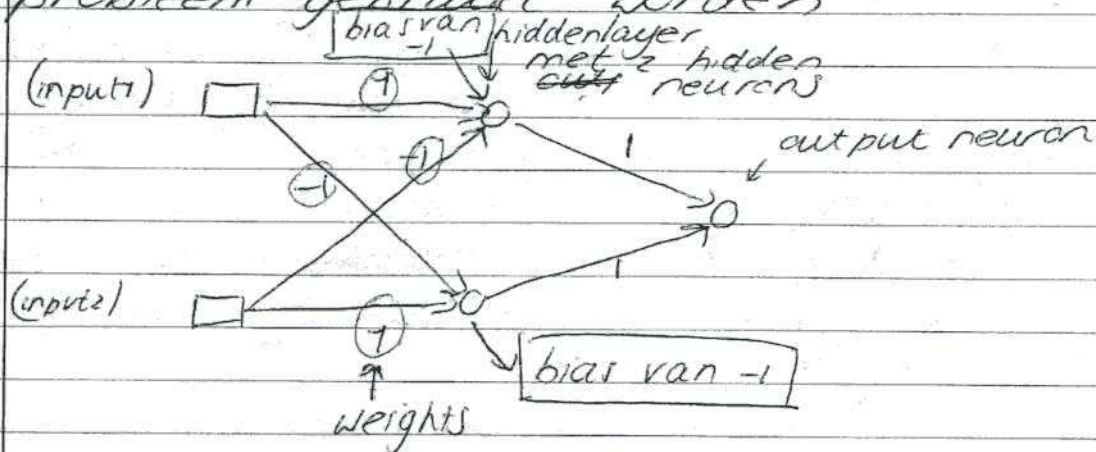
10  In dit geval moet gelden:

Input	gewenste output
(1, 1)	-1
(-1, -1)	-1
(1, -1)	1
(-1, 1)	1

We moeten dus een FFNN vinden die deze inputs als volgt scheidt: (d.m.v. 2 rechte lijnen)



Een FFNN die het XOR probleem oplost
 Een FFNN die het XOR probleem oplost
 bestaat uit twee input neuronen, twee
 hidden neuronen en 1 output neuron. Voor
 het neuron model van de hidden neuronen en
 de output neuron ~~geldt~~ geldt dat de sign
 activation function wordt gebruikt.
 Het volgende netwerk kan voor het XOR
 probleem gebruikt worden



Hierbij ga ik ervan uit dat de sign
 activation function als volgt werkt:

$$f(v_j) = \begin{cases} 1 & \text{als } v_j \geq 0 \\ -1 & \text{als } v_j < 0 \end{cases}$$

Met $v_j = \sum x_i \cdot w_{ij}$ $i=1,2$ $j=1,2$ (in dit geval)

* Voorbeeld (1, -1)

Output eerste hidden neuron =

$$\text{sign}(1 \cdot 1 + -1 \cdot -1 \text{ (bias)}) = \text{sign}(1) = 1$$

* Output tweede hidden neuron =

$$\text{sign}(1 \cdot -1 + -1 \cdot 1 \text{ (bias)}) = \text{sign}(-3) = -1$$

Output output neuron =

$$\text{sign}(1 \cdot 1 - 1 \cdot -1) = \text{sign}(2) = 1 \quad (\text{klopt!})$$



Faculteit der Exacte Wetenschappen

(Duidelijk en met blokletters invullen)

Naam en voorletters:

Cijfer:

Vak: *Neurale Netwerken*

Studentnummer: ,

Datum:

26-05-03

Jaar van 1e inschrijving: *99*

Studierichting:

2 J.

5

• Een radial basis function (RBF) is een functie waarvan de output afhangt van een niet stijgende ^{functie} ~~functie~~ afstand tussen een input vector en een stored vectors.

Bij een radial basis function wordt dus de output van een input vector berekend a.d.h.v. 'stored vectors' en de gewichten van deze ~~stored network vectors~~ ~~stored vectors~~ ^{betreffende} bij deze stored vectors.

de afstand van deze input tot stored vectors

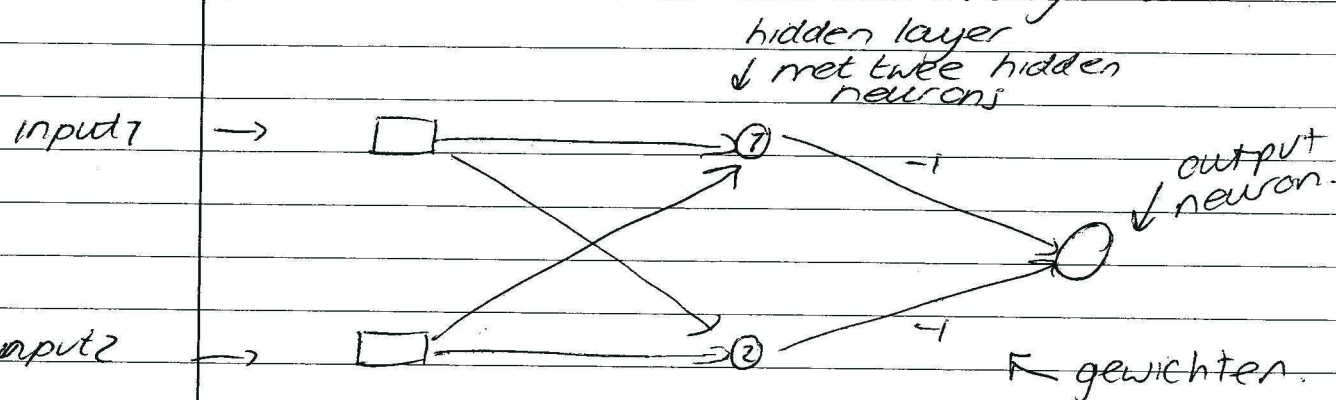
• Een RBF network bestaat ~~dus~~ uit $M+1$

10

hidden neurons die ieder een 'eigen' RBF activation function hebben en een 'eigen' center t_i ($i=1, \dots, M$). De input van de activation function van een bepaald hidden neuron is (in ieder geval)

$\|x - t_i\| \Rightarrow$ de afstand tussen het input exemplen ~~de~~ de center t_i die bij die ~~neuron~~ hidden neuron hoort. Er zijn verschillende mogelijkheden voor RBF activation functions. In bepaalde gevallen wordt de gaussian RBF activation function gebruikt, dan hebben we ook de spreiding (σ) nodig van de centers

10. Een RBF netwerk die het XOR probleem oplost bestaat uit 2 input neuronen, 2 hidden neuronen (met RBF activation functions) (met RBF gaussian activation functions.) Het netwerk ziet er als volgt uit:



Met center dat bij hidden neuron 1 hoort is $t_1 = (1, 1)$. Met center dat bij hidden neuron 2 hoort is $t_2 = (0, 0)$. De RBF gaussian activation die bij hidden neuron 1 hoort is $\varphi_1 = e^{-||x - t_1||^2}$ met $t_1 = (1, 1)$. De RBF gaussian activation function die bij hidden neuron 2 hoort is $\varphi_2 = e^{-||x - t_2||^2}$ met $t_2 = (0, 0)$.

De output van het netwerk geeft dat de input tot class 1 behoort als $-e^{-||x - t_1||^2} - e^{-||x - t_2||^2} + 1 > 0$, anders behoort de input tot class -1.

$x_1 = (1, 1)$
 $x_2 = (0, 0)$

$$\begin{pmatrix} 1 \\ 1 \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} = -1 - e^{-2} = -e^{-4} - e^{-2}$$

$$\begin{pmatrix} 0 \\ 0 \end{pmatrix} \begin{pmatrix} 1 \\ -1 \end{pmatrix} = -e^{-2} - 1$$

$$\begin{pmatrix} 1 \\ 1 \end{pmatrix} \begin{pmatrix} 1 \\ 1 \end{pmatrix} = 2$$

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

4. Het K-means algoritme:

- 10 - initialiseer de gewichten w random. Voor ieder gewicht representeert een cluster.
- Kijk voor iedere input bij welke cluster deze hoort door naar de afstand tussen deze input en de gewichten te kijken. Kies voor het gewicht met de minimale afstand ~~van het input example~~, tussen dit gewicht en het input example:

$$c(x) = \underset{j}{\operatorname{argmax}} \|x - w_j(n)\|$$

- Bereken de gewichten nu als volgt:

$$w_j(n+1) = \frac{\sum x_j}{n_j}$$

↓
voor alle x die tot het cluster van w_j behoren.

n_j = aantal x gewichten van in cluster j .

n_j = aantal x (example) in cluster van w_j .

- Na nog n op: $n = n+1$, en ga door totdat de gewichten niet tot bijna niet meer veranderen.

5. De gewichten van het Hopfield netwerk worden als volgt berekend:

(N = input dimensie)

(M = ~~aantal de~~ aantal 'stored' memorie vectoren')

$$w_{ij} = \frac{1}{M} \sum_{v=1}^M f(v_i) f(v_j) \quad \text{als } i \neq j$$

$$w_{ij} = \frac{1}{M} \sum_{v=1}^M f(v_i) f(v_j) \quad \text{als } i=j$$

$f(v_i)$ = ^{waarde van} i component van v stored memorie vector

Dit houdt in dat als componenten i en j in veel stored memories bijvoorbeeld positief gecorreleerd zijn dat w_{ij} een groot getal is.
 waarschijnlijk

Dus: w_{ij} hangt af van de correlatie afhankelijkheid tussen componenten i en j in stored memorie vectoren.

10 • Compare Hopfield en BSB networks:

* Hopfield en BSB networks zijn beide vormen van associated memories networks. De networks komen ~~ook~~ op verschillende punten overeen:

- bij beide networks daalt, iedere keer bij een verandering van de neuron states, de energie functie.
- Dit houdt in dat de energie dus daalt, en uiteindelijk, als 'de attractor state bereikt is, de energie tot een minimum is gedaald.
- Beide networks werken met attractor states.
- Beide maken gebruik van feed back.
Verschil: in een Hopfield network is er geen sprake van self feed back, in een BSB wel.
- Beide networks maken gebruik van 'Hebb's postulate of learning'.

Verdere verschillen zitten m.n. in de toepassingen:

- BSB staat bekend om zijn cluster-vaardigheden zoals concept formation
- Hopfield networks zijn meer gericht op 'content addressable memories' en optimalisatie.

Andere verschil: Bij Hopfield wordt de state vector x aangepast, bij BSB wordt in iedere iteratie de gehele state vector aangepast. (Bij Hopfield wordt één component van de state vector aangepast, bij BSB wordt de gehele state vector aangepast.)