

Aprendizado Auto Supervisionado

Representação por Componentes Principais

PCA - Principal Component Analysis

- Análise de Componentes Principais
- Transformada de Kahunen-Löeve

via Redes Neurais

$$\underline{x} \Rightarrow \underline{z} \Rightarrow \underline{\tilde{x}} \cong \underline{x}$$

$$\underline{x} \Rightarrow \underline{z} \Rightarrow \underline{\tilde{x}} \cong \underline{x}$$

Critérios e Propriedades da PCA

**Redução de dimensionalidade,
Compressão da informação**

$$\dim \underline{z} < \dim \underline{x}$$

Mapeamentos Lineares

$$\underline{z} = \underline{W} \underline{x} \quad \underline{\tilde{x}} = \underline{T} \underline{z}$$

Representação ótima

$$\underline{W}, \underline{T} = \arg \min E(|\underline{x} - \underline{\tilde{x}}|^2)$$

Base Ortonormal

$$\underline{W}\underline{W}^t = \underline{I}$$

Descorrelação das componentes

$$E(z_i z_j) = 0 \quad \forall i, j$$

**Ordem de importância,
Potência das componentes**

$$E(z_i^2) > E(z_j^2) \quad \forall i < j$$

consequência:

$$\underline{z} = \underline{W} \underline{x} \quad \underline{\tilde{x}} = \underline{T} \underline{z} = \underline{T} \underline{W} \underline{x}$$

se $\underline{\tilde{x}} \cong \underline{x}$ então $\underline{T} \underline{W} \cong \underline{I}$ e $\underline{T} \cong \underline{W}^{-1}$ *pseudo-inversa*

mas se $\underline{W}$ é uma base ortonormal

então $\underline{W}^{-1} = \underline{W}^t$ logo $\underline{T} \cong \underline{W}^t$

$$\text{e} \quad \underline{z} = \underline{W} \underline{x} \quad \underline{\tilde{x}} = \underline{W}^t \underline{z}$$

Cálculo da PCA (W) via Redes Neurais

Pré processamento

$$E(\underline{x}) = \underline{0} \quad \Rightarrow \quad E(\underline{z}) = E(\underline{\tilde{x}}) = \underline{0}$$

e de preferência $\sigma(x_i) \cong 1 \quad \forall i$

Implementação dos mapeamentos

$$\underline{z} = \underline{W} \underline{x}$$

$$\dim \underline{x} = n \quad \dim \underline{z} = m \leq n$$

$$\underline{W} \text{ é } m \times n$$

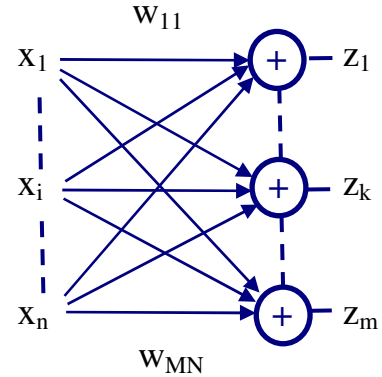
$$\underline{W} = \begin{bmatrix} w_{11} & w_{12} & \dots & w_{1n} \\ w_{21} & w_{22} & \dots & \\ \dots & & & \\ w_{m1} & \dots & & w_{mn} \end{bmatrix} = \begin{bmatrix} \underline{w}_1^t \\ \underline{w}_2^t \\ \dots \\ \underline{w}_m^t \end{bmatrix}$$

$$z_i = \underline{w}_i^t \underline{x} = \sum_{j=1}^n w_{ij} x_j$$

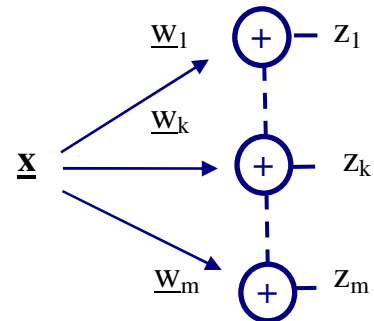
$$\underline{z} = \sum_{j=1}^n x_j \underline{w}_j$$

$$\underline{z} = \underline{W} \underline{x}$$

$$z_i = \underline{w}_i^t \underline{x}$$



**In-star
de
Grossberg**



$$\tilde{\underline{x}} = \underline{T} \underline{z} = \underline{W}^t \underline{z}$$

$$\dim \tilde{\underline{x}} = n \qquad \dim \underline{z} = m \leq n$$

$$\underline{T} = \underline{W}^t = \begin{bmatrix} w_{11} & w_{21} & \dots & w_{m1} \\ w_{12} & w_{22} & \dots & \\ \dots & & & \\ w_{1n} & \dots & & w_{mn} \end{bmatrix}$$

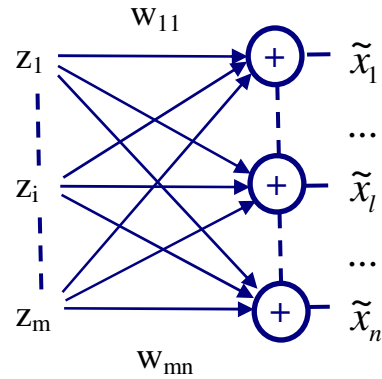
$$= \begin{bmatrix} \underline{w}_1 & \underline{w}_2 & \dots & \underline{w}_m \end{bmatrix}$$

$$\tilde{x}_i = \sum_{j=1}^m w_{ji} z_j$$

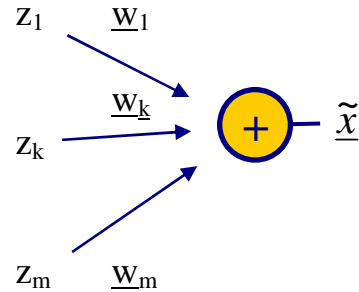
$$\tilde{\underline{x}} = \sum_{j=1}^m z_j \underline{w}_j$$

$$\underline{\tilde{x}} = \underline{W}^t \underline{z}$$

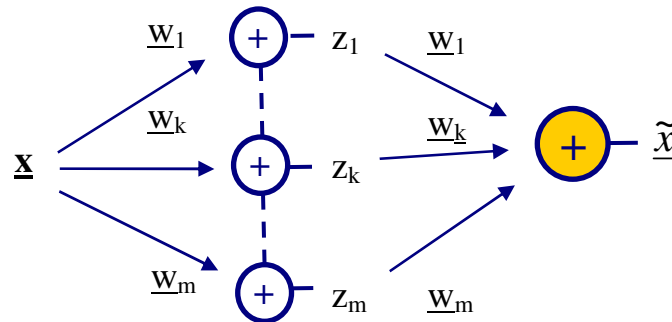
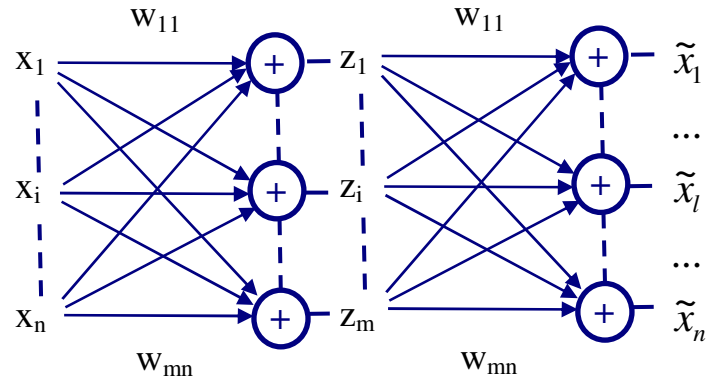
$$\underline{\tilde{x}} = \sum_{j=1}^m z_j \underline{w}_j$$



**Out-star
de
Grossberg**

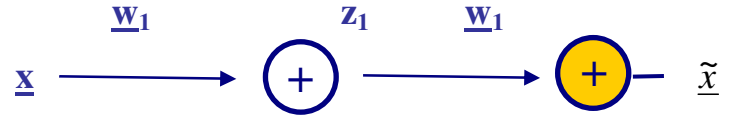


Mapeamento completo



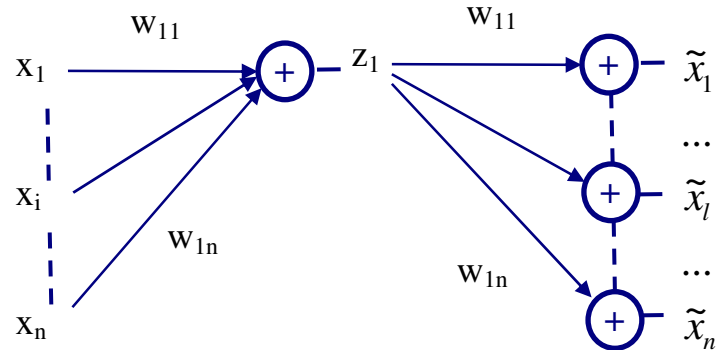
Primeira componente da base, $\underline{w}_1$

$$z_1 = \underline{w}_1^t \underline{x} = \sum_{j=1}^n w_{1j} x_j$$



$$\tilde{x} = z_1 \underline{w}_1 = \underline{w}_1^t \underline{x} \underline{w}_1$$

$$\tilde{x}_i = w_{1i} z_1 = w_{1i} \sum_{j=1}^n w_{1j} x_j$$



Cálculo da primeira componente, $\underline{w}_1$

Critério

$$\underline{w}_1 = \arg \text{Min } E(|\underline{x} - \tilde{\underline{x}}|^2)$$

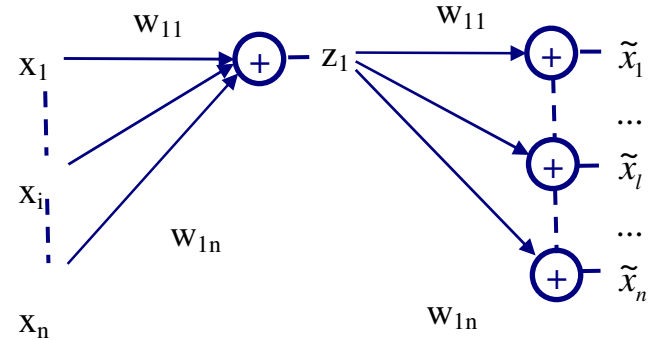
Método : Gradiente descendente

$$\Delta w_{1k} = -\alpha \frac{\partial}{\partial w_{1k}} E(|\underline{x} - \tilde{\underline{x}}|^2)$$

$$|\underline{x} - \underline{\tilde{x}}|^2 = \mathcal{E}^2 = \sum_{i=1}^n \mathcal{E}_i^2 = \sum_{i=1}^n (x_i - \tilde{x}_i)^2$$

$$\tilde{x}_i = z_1 w_{1i}$$

$$z_1 = \sum_{j=1}^n w_{1j} x_j$$



donde :

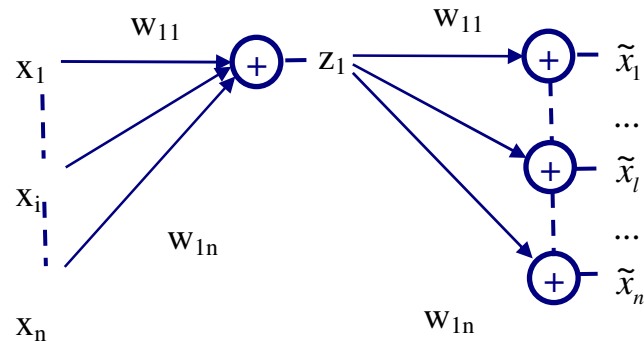
$$\Delta w_{1k} = 2\alpha E \left(z_1 \mathcal{E}_k + x_k \sum_{i=1}^n \mathcal{E}_i w_{1i} \right)$$

Treinamento da Primeira Componente, $\underline{w}_1$

Inicialização para a primeira componente:

$$\underline{w}_1 \text{ inicial} = \underline{r} \quad \text{randômico,} \quad |\underline{w}_1 \text{ inicial}| \approx 1$$

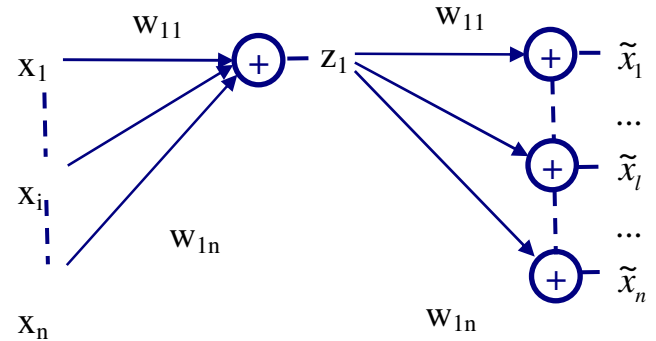
os valores iniciais das sinapses
equivalentes da
primeira e segunda camadas
são obrigatoriamente iguais



Treinamento da Primeira Componente $\underline{w}_1$

Para cada entrada $\underline{x}$

Signal feedforward,
saída,
erro



Acréscimos nas sinapses

$$\Delta w_{1k} = 2\alpha E \left(z_1 \varepsilon_k + x_k \sum_{i=1}^n \varepsilon_i w_{1i} \right)$$

Segunda componente, $\underline{w}_2$

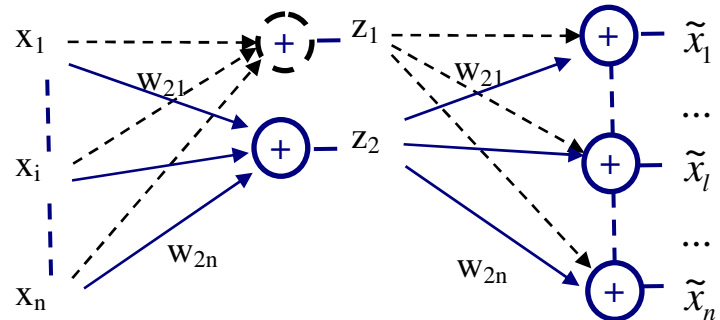
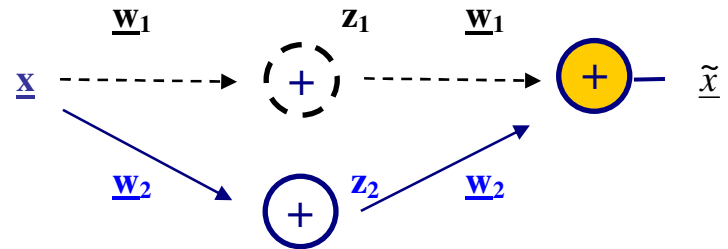
$$z_2 = \underline{w}_2^t \underline{x} = \sum_{j=1}^n w_{2j} x_j$$

$$\tilde{x} = z_1 w_1 + z_2 w_2$$

$$= z_1 w_1 + \underline{w}_2^t \underline{x} w_2$$

$$\tilde{x}_i = w_{1i} z_1 + w_{2i} z_2$$

$$= w_{1i} z_1 + w_{2i} \sum_{j=1}^n w_{2j} x_j$$



Cálculo da segunda componente, $\underline{w}_2$

Critério

$$\underline{w}_2 = \arg \min E(|\underline{x} - \tilde{x}|^2) \quad \text{com} \quad \underline{w}_1 = \underline{w}_{1\text{ótimo}}$$

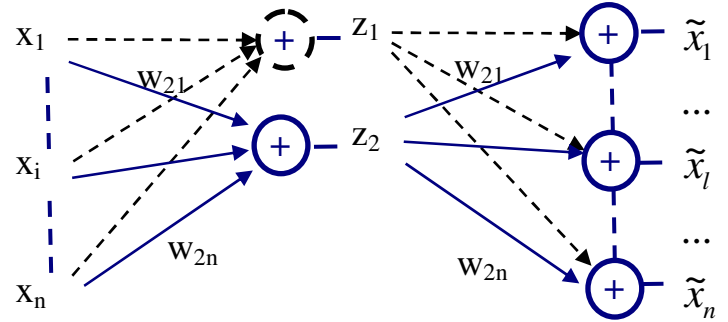
Método : Gradiente descendente

$$\Delta w_{2k} = -\alpha \frac{\partial}{\partial w_{2k}} E(|\underline{x} - \tilde{x}|^2)$$

$$|\underline{x} - \underline{\tilde{x}}|^2 = \mathcal{E}^2 = \sum_{i=1}^n \mathcal{E}_i^2 = \sum_{i=1}^n (x_i - \tilde{x}_i)^2$$

$$\tilde{x}_i = z_1 w_{1i} + z_2 w_{2i}$$

$$z_2 = \sum_{j=1}^n w_{2j} x_j$$



donde :

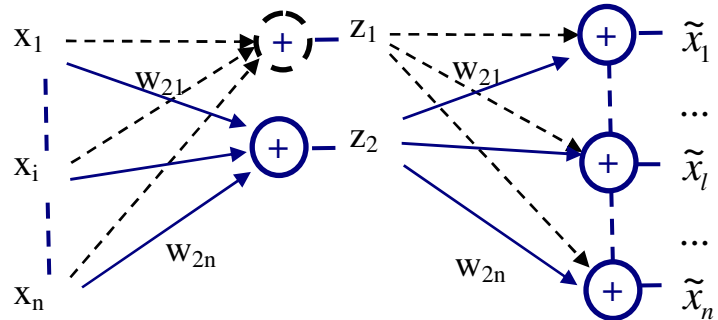
$$\Delta w_{2k} = 2\alpha E \left(z_2 \mathcal{E}_k + x_k \sum_{i=1}^n \mathcal{E}_i w_{2i} \right)$$

Treinamento da Segunda Componente

Inicialização para a segunda componente:

$$\underline{w}_2 \text{ inicial} = \underline{r} \quad \text{randômico}, \quad |\underline{w}_2 \text{ inicial}| \approx 1$$

os valores iniciais das sinapses
equivalentes da
primeira e segunda camadas
são obrigatoriamente iguais



Treinamento da Segunda Componente

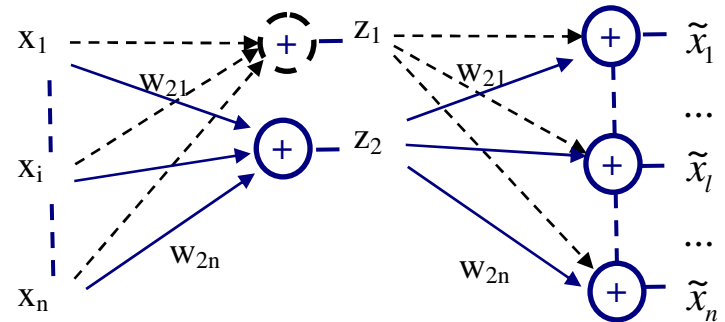
Congelar as sinapses anteriores w_1

Para cada entrada $\underline{x}$

Signal feedforward,
saída,
erro

Acréscimos nas sinapses

$$\Delta w_{2k} = 2\alpha E \left(z_2 \varepsilon_k + x_k \sum_{i=1}^n \varepsilon_i w_{2i} \right)$$



Terceira componente $\underline{w}_3$

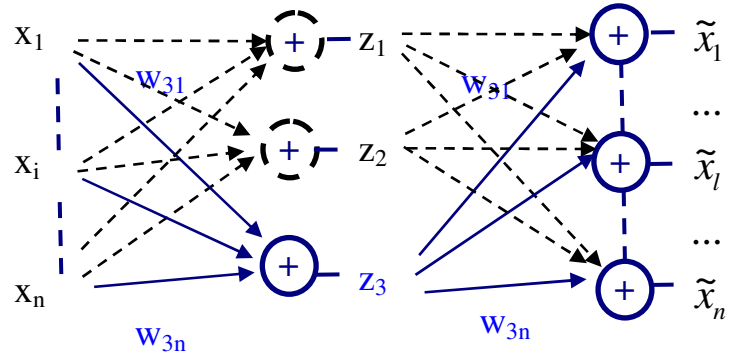
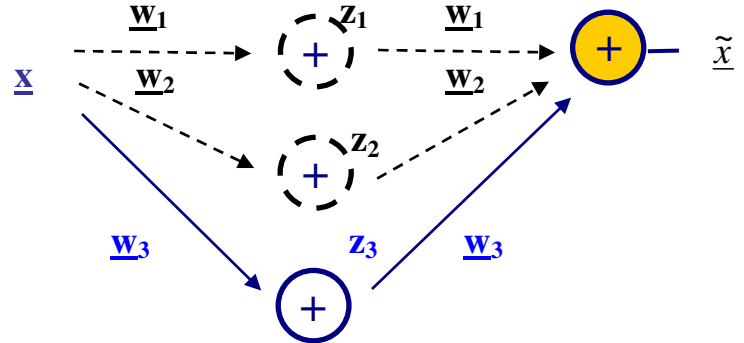
$$z_3 = \underline{w}_3^t \underline{x} = \sum_{j=1}^n w_{3j} x_j$$

$$\tilde{x} = z_1 \underline{w}_1 + z_2 \underline{w}_2 + z_3 \underline{w}_3$$

$$= z_1 \underline{w}_1 + z_2 \underline{w}_2 + \underline{w}_3^t \underline{x} \underline{w}_3$$

$$\tilde{x}_i = w_{1i} z_1 + w_{2i} z_2 + w_{3i} z_3$$

$$= w_{1i} z_1 + w_{2i} z_2 + w_{3i} \sum_{j=1}^n w_{3j} x_j$$



Cálculo da terceira componente $\underline{w}_3$

Crítério

$$\underline{w}_3 = \arg \text{Min } E(|\underline{x} - \tilde{\underline{x}}|^2)$$

com $\underline{w}_1 = \underline{w}_{1\acute{o}timo}$ e $\underline{w}_2 = \underline{w}_{2\acute{o}timo}$

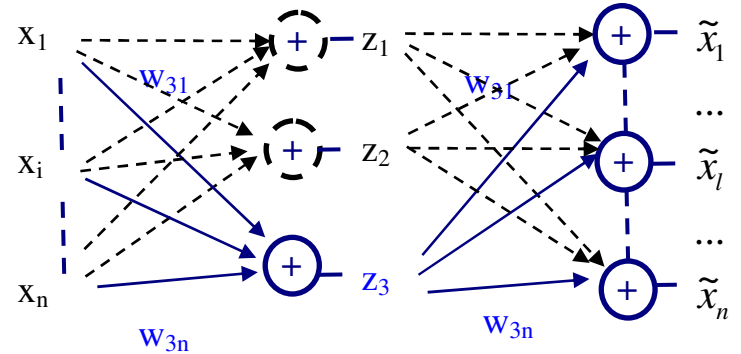
Método : Gradiente descendente

$$\Delta w_{3k} = -\alpha \frac{\partial}{\partial w_{3k}} E(|\underline{x} - \tilde{\underline{x}}|^2)$$

$$|\underline{x} - \underline{\tilde{x}}|^2 = \mathcal{E}^2 = \sum_{i=1}^n \mathcal{E}_i^2 = \sum_{i=1}^n (x_i - \tilde{x}_i)^2$$

$$\tilde{x}_i = z_1 w_{1i} + z_2 w_{2i} + z_3 w_{3i}$$

$$z_3 = \sum_{j=1}^n w_{3j} x_j$$



donde :

$$\Delta w_{3k} = 2\alpha E \left(z_3 \mathcal{E}_k + x_k \sum_{i=1}^n \mathcal{E}_i w_{3i} \right)$$

Treinamento da Terceira Componente $\underline{w}_3$

Congelar as sinapses anteriores $\underline{w}_1$ e $\underline{w}_2$

Inicialização para a terceira componente $\underline{w}_3$:

$$\underline{w}_3 \text{ inicial} = \underline{r} \quad \text{randômico,} \quad |\underline{w}_3 \text{ inicial}| \approx 1$$

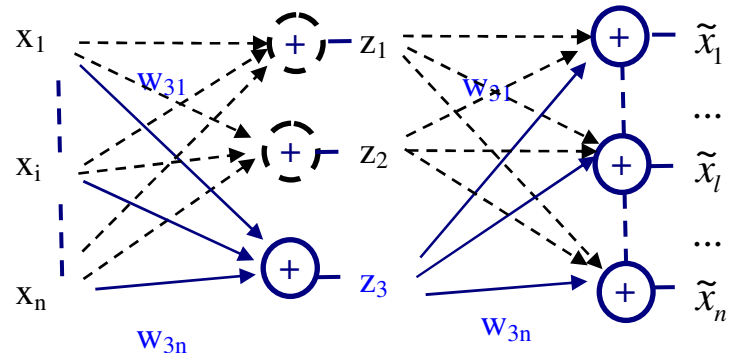
Para cada entrada $\underline{x}$

Signal feedforward,

saída,

erro

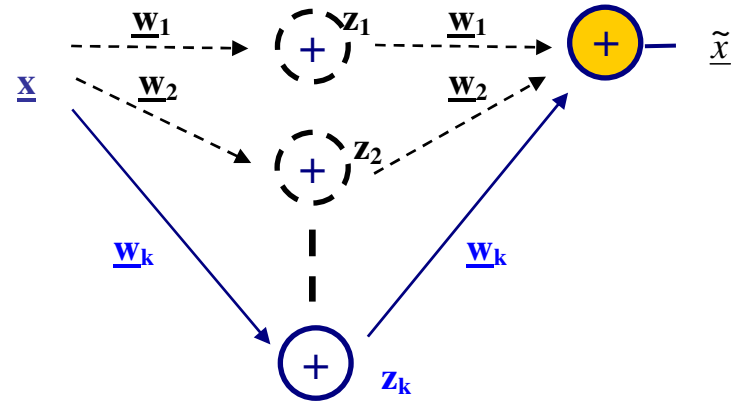
Acréscimos sinapses



$$\Delta w_{3k} = 2\alpha E \left(z_3 \mathcal{E}_k + x_k \sum_{i=1}^n \mathcal{E}_i w_{3i} \right)$$

k-ésima Componente

– análogo



Observação sobre a ortogonalidade das componentes:

Pelo processo de cálculo as componentes são naturalmente ortogonais.

utilizando apenas a primeira componente

$$\tilde{\underline{x}}_1 = \underline{w}_1^t \underline{x} \underline{w}_1 \qquad \underline{w}_1 = \arg \operatorname{Min} E(|\underline{x} - \tilde{\underline{x}}|^2)$$

$$\underline{\varepsilon}_1 = \underline{x} - \tilde{\underline{x}}_1 = \underline{x} - \underline{w}_1^t \underline{x} \underline{w}_1 \qquad \underline{w}_1^t \cdot \underline{\varepsilon}_1 = 0$$

utilizando as duas primeiras componentes

$$\underline{\tilde{x}}_2 = \underline{w}_1^t \underline{x} \underline{w}_1 + \underline{w}_2^t \underline{x} \underline{w}_2 \quad \underline{w}_2 = \arg \Big|_{\underline{w}_1 = \underline{w}_1 \text{ ótimo}} \text{Min } E \left(\left| \underline{x} - \underline{\tilde{x}} \right|^2 \right)$$

$$\underline{\mathcal{E}}_2 = \underline{x} - \underline{\tilde{x}}_2 = \underline{x} - \underline{w}_1^t \underline{x} \underline{w}_1 - \underline{w}_2^t \underline{x} \underline{w}_2 = \underline{\mathcal{E}}_1 - \underline{w}_2^t \underline{x} \underline{w}_2$$

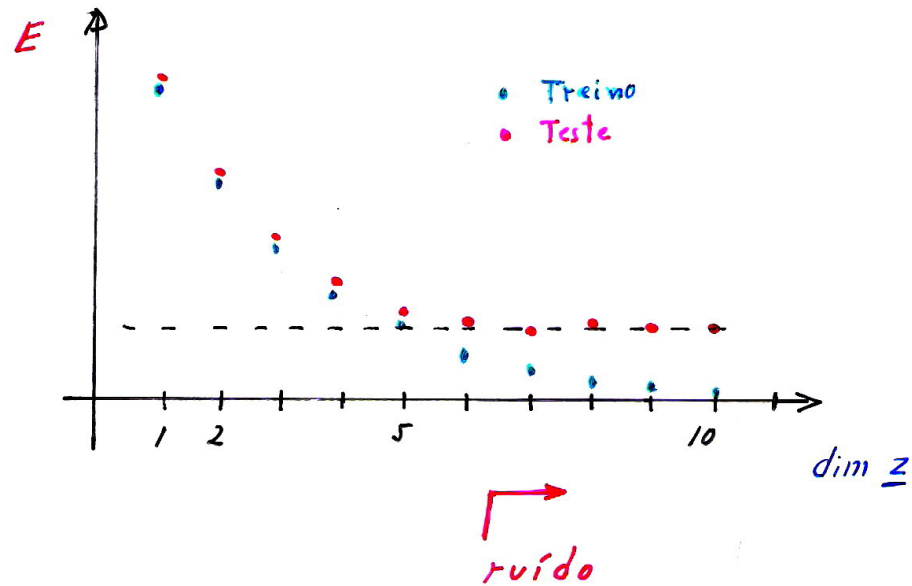
$$\underline{w}_2^t \cdot \underline{\mathcal{E}}_2 = 0 \quad \underline{w}_1^t \cdot \underline{\mathcal{E}}_2 = 0$$

$$\underline{w}_1^t \cdot \underline{\mathcal{E}}_2 = \underline{w}_1^t \left(\underline{\mathcal{E}}_1 - \underline{w}_2^t \underline{x} \underline{w}_2 \right) = \underline{w}_1^t \underline{\mathcal{E}}_1 - \underline{w}_1^t \left(\underline{w}_2^t \underline{x} \right) \underline{w}_2 = 0 - \left(\underline{w}_2^t \underline{x} \right) \underline{w}_1^t \underline{w}_2 = 0$$

$$\text{donde} \quad \underline{w}_1^t \underline{w}_2 = 0$$

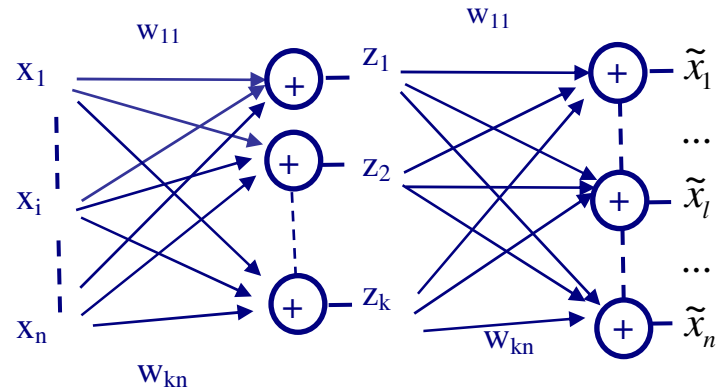
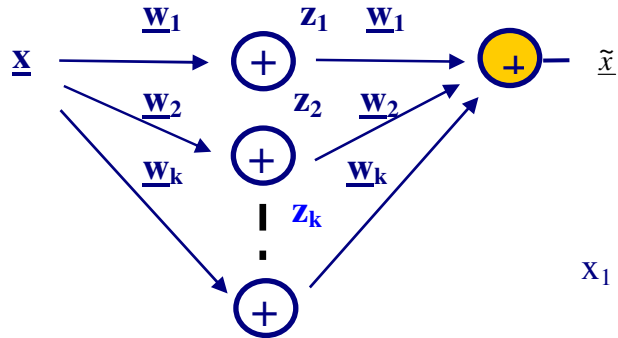
Erro $E(|\underline{x} - \tilde{\underline{x}}|^2)$ vs.

número de componentes = $k = \dim \underline{z}$



Processo alternativo:

Treinar todas as k componentes simultaneamente



Cálculo das componente, $\underline{w}_1, \dots, \underline{w}_k$

Cr  rio

$$\{\underline{w}_1, \dots, \underline{w}_k\} = \arg \min E(|\underline{x} - \tilde{\underline{x}}|^2)$$

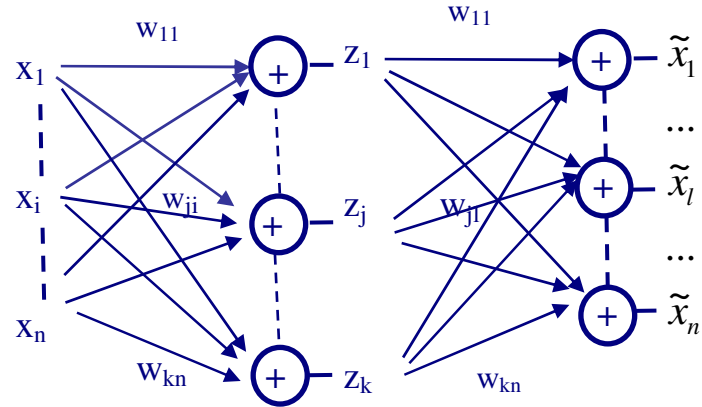
M  todo : Gradiente descendente

$$\Delta w_{ij} = -\alpha \frac{\partial}{\partial w_{ij}} E(|\underline{x} - \tilde{\underline{x}}|^2)$$

$$|\underline{x} - \underline{\tilde{x}}|^2 = \mathcal{E}^2 = \sum_{l=1}^n \mathcal{E}_l^2 = \sum_{l=1}^n (x_l - \tilde{x}_l)^2$$

$$\tilde{x}_l = \sum_{j=1}^k z_j w_{lj}$$

$$z_j = \sum_{i=1}^n w_{ji} x_i$$



donde :

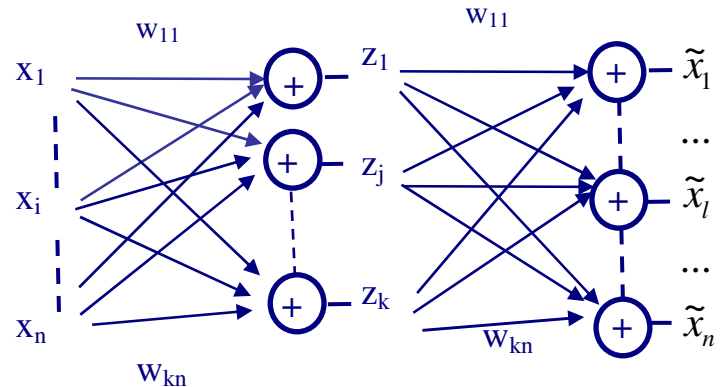
$$\Delta w_{ij} = 2\alpha E \sum_{i=1}^n \left(z_i \mathcal{E}_j + x_j \sum_{k=1}^n \mathcal{E}_k w_{ik} \right) \quad (\text{\textcolor{brown}{à verificar}})$$

Treinamento da Rede

Inicialização para cada componente $\underline{w}_k$

$$\underline{w}_{k \text{ inicial}} = \underline{r} \quad \text{randômico,} \quad |\underline{w}_{k \text{ inicial}}| \approx 1$$

os valores iniciais das sinapses
equivalentes da
primeira e segunda camadas
são obrigatoriamente iguais



Treinamento das Componentes $\underline{w}_1, \dots, \underline{w}_k$

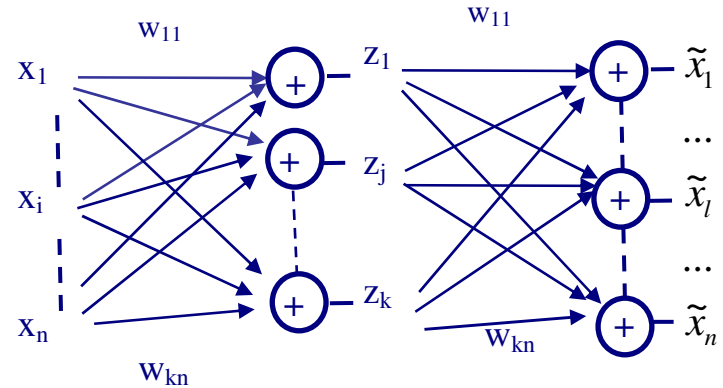
Para cada entrada $\underline{x}$

Signal feedforward,

saída,

erro

Acréscimos nas sinapses



$$\Delta w_{ij} = 2\alpha E \sum_{i=1}^n \left(z_i \mathcal{E}_j + x_j \sum_{k=1}^n \mathcal{E}_k w_{ik} \right) \quad (\text{\textcolor{brown}{à verificar}})$$

Restrição

$$\{\underline{w}_1, \dots, \underline{w}_k\}$$

obtido é apenas uma base (**não ortonormal, não ordenada**) para o espaço das k primeiras componentes principais.