

Treinamento da Rede Neural =

Minimização de F_{in}

Otimização (minimização) de Funções Contínuas

$$F(w_1, w_2, \dots, w_n) = F(\underline{w})$$

$$\underline{w} = \begin{bmatrix} w_1 \\ w_2 \\ \vdots \\ w_n \end{bmatrix}$$

Propriedades

$$F(\underline{w}) \geq 0$$

$$\frac{\partial F}{\partial w_i} \quad \exists \quad \forall w_i, \quad \forall \underline{w}$$

Como encontrar $\underline{w}^*$ tal que

$$F(\underline{w}^*) \leq F(\underline{w}) \quad \forall \underline{w} \neq \underline{w}^* \quad ?$$

Otimização de F

não linear

não convexa

w ordem elevada

irrestrita

contínua, derivável

Como fazer ???

Noções Elementares de Otimização

1 – Função objetivo à minimizar

(encontrar o ponto de mínimo ou minimante)

$$F(w_1, w_2, \dots) = F(\underline{w})$$

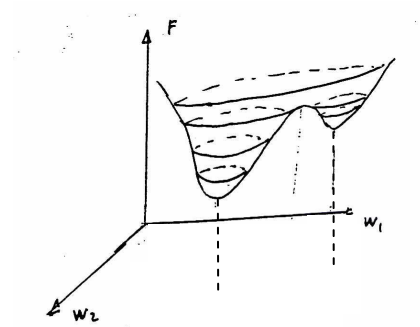
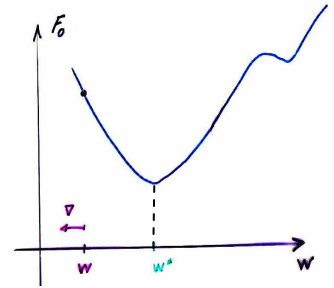
Mínimo, minimante $\underline{w}^*$

$$F(\underline{w}^*) < F(\underline{w}^* + \underline{\Delta w}) \quad \forall \underline{\Delta w} \mid |\underline{\Delta w}| < \varepsilon$$

→ mínimo local

Mínimo global

$$F(\underline{w}^*) \leq F(w) \quad \forall w \neq \underline{w}^*$$



2 – Aproximações no entorno de um ponto qualquer $\underline{w}_0$ (Taylor)

$$\underline{w} = \underline{w}_0 + \underline{\Delta w} \quad |\underline{\Delta w}| \leq \varepsilon$$

$$F(\underline{w}_0 + \underline{\Delta w}) = F(\underline{w}_0) + \sum_i \left. \frac{\partial F}{\partial w_i} \right|_{\underline{w}_0} \Delta w_i + \frac{1}{2} \sum_i \sum_j \left. \frac{\partial^2 F}{\partial w_i \partial w_j} \right|_{\underline{w}_0} \Delta w_i \Delta w_j + \dots$$

Aproximação linear ou de 1ª. ordem

$$F(\underline{w}_0 + \underline{\Delta w}) \approx F(\underline{w}_0) + \sum_i \left. \frac{\partial F}{\partial w_i} \right|_{\underline{w}_0} \Delta w_i$$

a aproximação é válida em um pequeno entorno de $\underline{w}_0$

Notação usando vetores

vetor **variáveis livres**

$$\underline{w} = \begin{bmatrix} w_1 \\ w_2 \\ \vdots \\ w_n \end{bmatrix}$$

vetor **acréscimos** nas variáveis livres

$$\underline{\Delta w} = \begin{bmatrix} \Delta w_1 \\ \Delta w_2 \\ \vdots \\ \Delta w_n \end{bmatrix}$$

vetor **gradiente**

$$\nabla F(\underline{w}_0) = \left[\frac{\partial F}{\partial w_i} \bigg|_{\underline{w}_0} \right] = \begin{bmatrix} \frac{\partial F}{\partial w_1} \bigg|_{\underline{w}_0} \\ \frac{\partial F}{\partial w_2} \bigg|_{\underline{w}_0} \\ \vdots \\ \frac{\partial F}{\partial w_n} \bigg|_{\underline{w}_0} \end{bmatrix}$$

$$F(\underline{w}_0 + \underline{\Delta w}) = F(\underline{w}_0) + \sum_i \frac{\partial F}{\partial w_i} \bigg|_{\underline{w}_0} \Delta w_i + \dots$$

$$F(\underline{w}_0 + \underline{\Delta w}) = F(\underline{w}_0) + \underline{\Delta w}^t \underline{\nabla F}(\underline{w}_0) + \dots$$

Aproximação linear ou de 1ª. ordem

$$F(\underline{w}_0 + \underline{\Delta w}) \simeq F(\underline{w}_0) + \underline{\Delta w}^t \underline{\nabla F}(\underline{w}_0)$$

válida em um pequeno entorno de $\underline{w}_0$

3 - Mínimo, minimante $\underline{w}^*$

$$F(\underline{w}^* + \underline{\Delta w}) = F(\underline{w}^*) + \Delta F(\underline{w}^*) > F(\underline{w}^*) \quad \forall \underline{\Delta w} \quad | \quad |\underline{\Delta w}| \leq \varepsilon$$

$$F(\underline{w}^* + \underline{\Delta w}) \cong F(\underline{w}^*) + \underline{\Delta w}^t \underline{\nabla F}$$

$$\Delta F = \underline{\Delta w}^t \underline{\nabla F} = |\underline{\Delta w}| |\underline{\nabla F}| \cos \angle \underline{\Delta w}, \underline{\nabla F} > 0$$

+ + +/-

$$\Rightarrow \underline{\nabla F}(\underline{w}^*) = \underline{0}$$

Nos pontos de mínimo (minimantes) o gradiente é nulo

4 - Propriedades do Gradiente

4.1 - Nos pontos de mínimo (minimantes) o gradiente é nulo

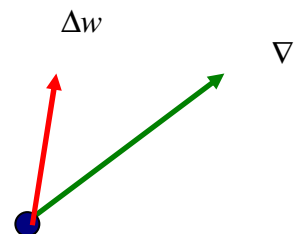
$$\underline{\nabla F}(\underline{w}^*) = \underline{0}$$

4.2 - Para $\underline{\Delta w}$ pequeno e constante o gradiente indica a direção de máximo ΔF (e $-\underline{\nabla F}$ indica a direção de maior decréscimo de F)

vale a aproximação de primeira ordem $|\underline{\Delta w}| \leq \varepsilon$

$$F(\underline{w}_0 + \underline{\Delta w}) = F(\underline{w}_0) + \Delta F \cong F(\underline{w}_0) + \underline{\Delta w}^t \underline{\nabla F}$$

$$\Delta F \cong \underline{\Delta w}^t \underline{\nabla F} = |\underline{\Delta w}| |\underline{\nabla F}| \cos \theta$$

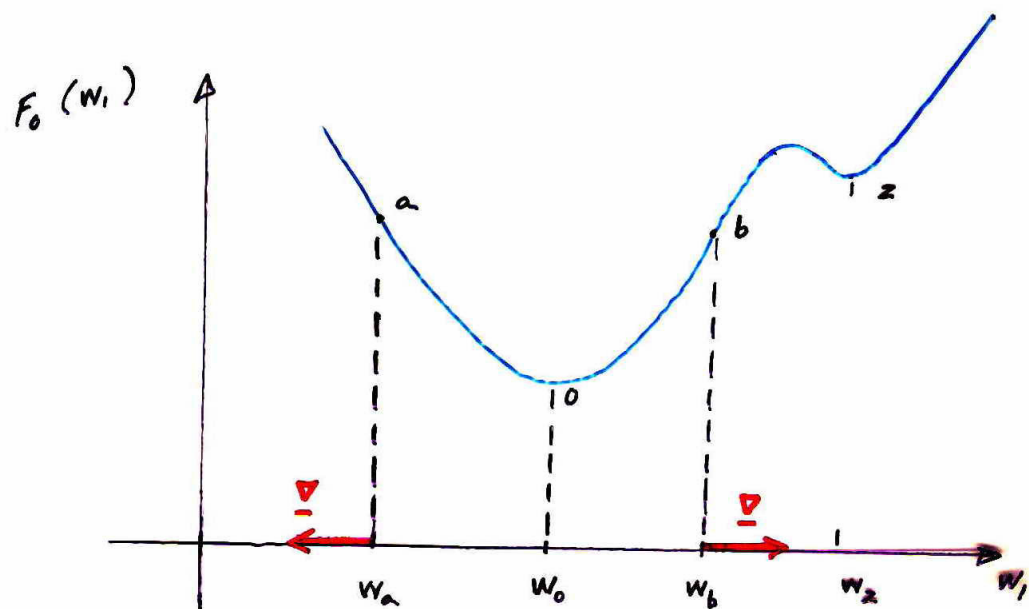
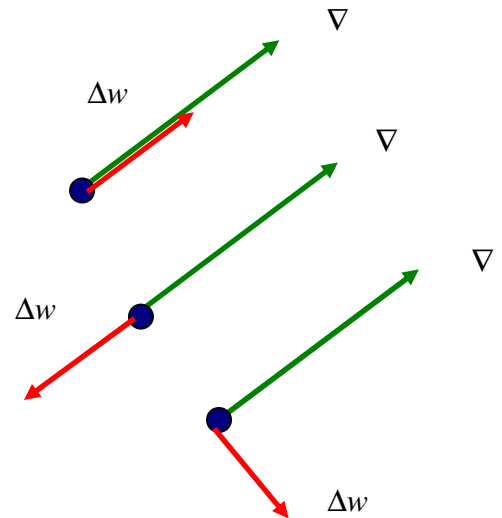


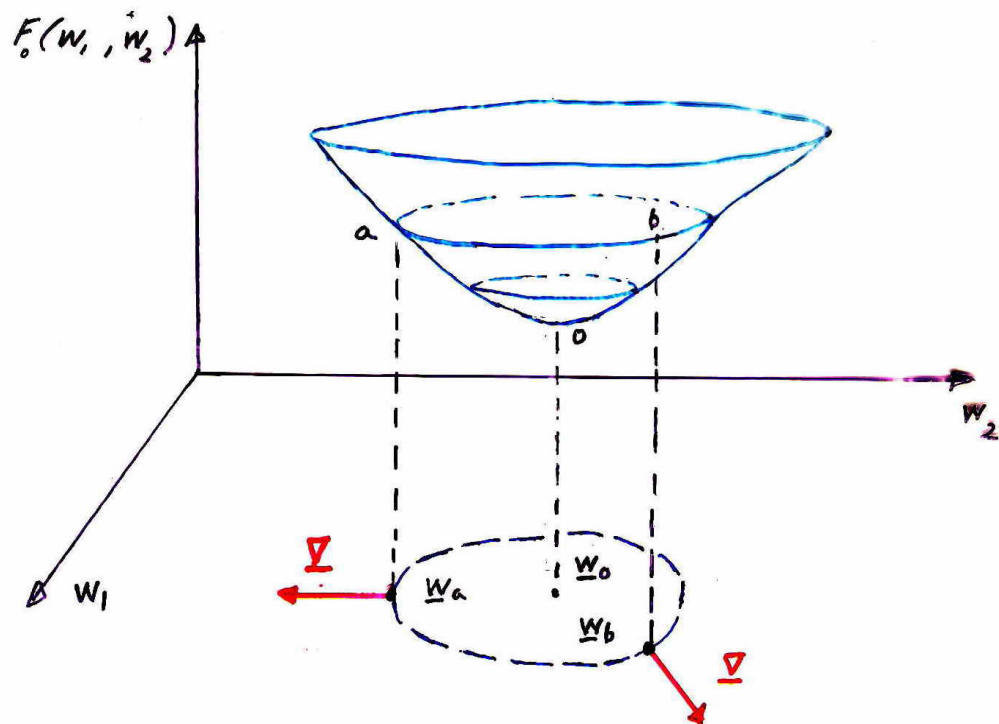
$$\Delta F \cong \underline{\Delta w}^t \nabla F = |\underline{\Delta w}| |\underline{\nabla F}| \cos \theta$$

$$\theta = 0 \quad \longrightarrow \quad \Delta F \text{ max}$$

$$\theta = \pi \quad \longrightarrow \quad \Delta F \text{ min} \\ (-\Delta F \text{ max})$$

$$\theta = \pm \frac{\pi}{2} \quad \longrightarrow \quad \Delta F = 0$$





Otimização

Determinação do minimante de $F(\underline{w})$

$\underline{w}^* \mid F(\underline{w}^*)$ é mínima

ou

$\underline{w}^* \mid \nabla F(\underline{w}^*) = \underline{0}$

$$\frac{\partial F}{\partial w_i} = 0 \quad \forall w_i$$

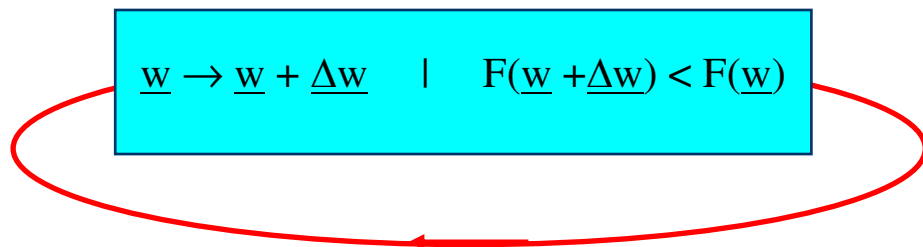
Otimização: Solução Analítica

$$\underline{w}^* \quad | \quad \nabla F(\underline{w}^*) = \underline{0}$$

usualmente muito complexo

Otimização: Métodos Numéricos Recursivos

Descida continuada



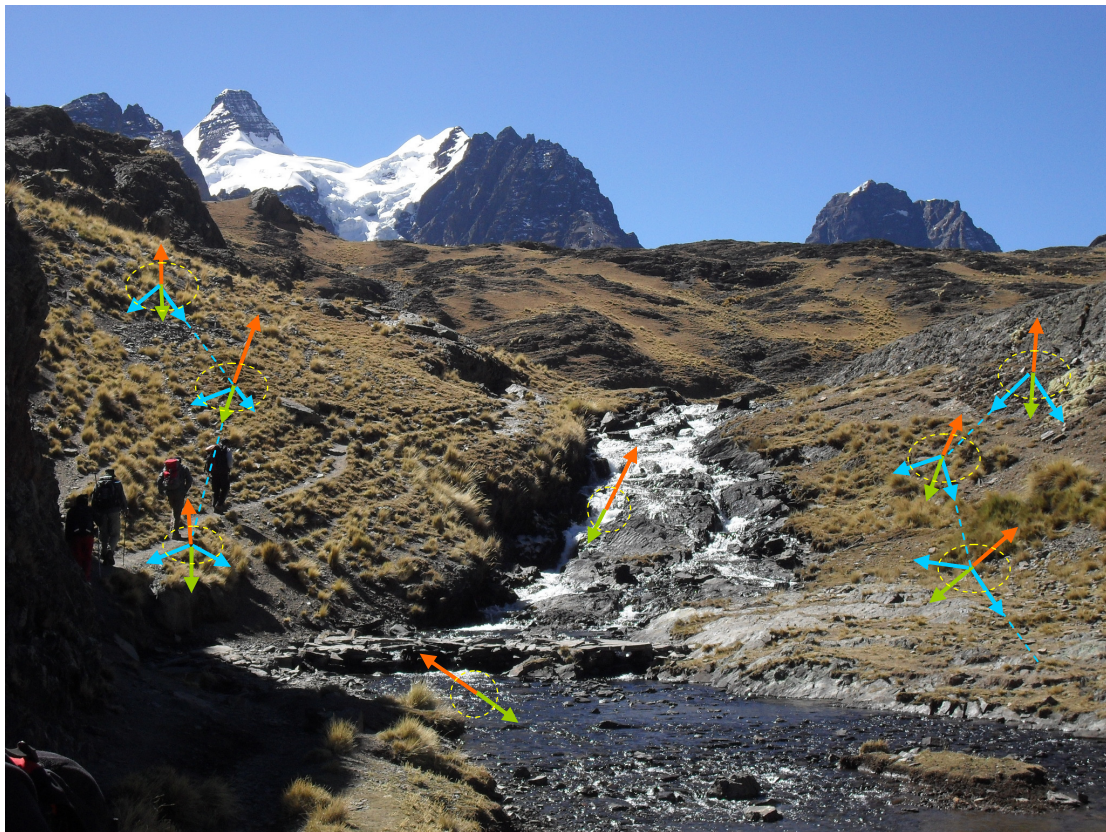
$\underline{\Delta w}$??

$$\underline{\Delta w} = \alpha \underline{d}$$

Passo
controla $|\underline{\Delta w}|$
 $\alpha > 0$, pequeno

controla
direção $\underline{\Delta w}$

Descida continuada, permanente - **decisão local**

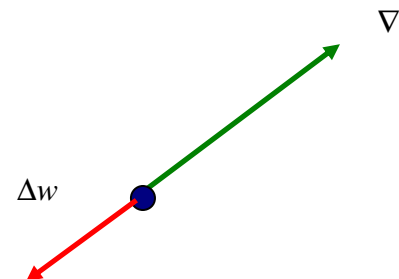


Método do Gradiente Descendente ou

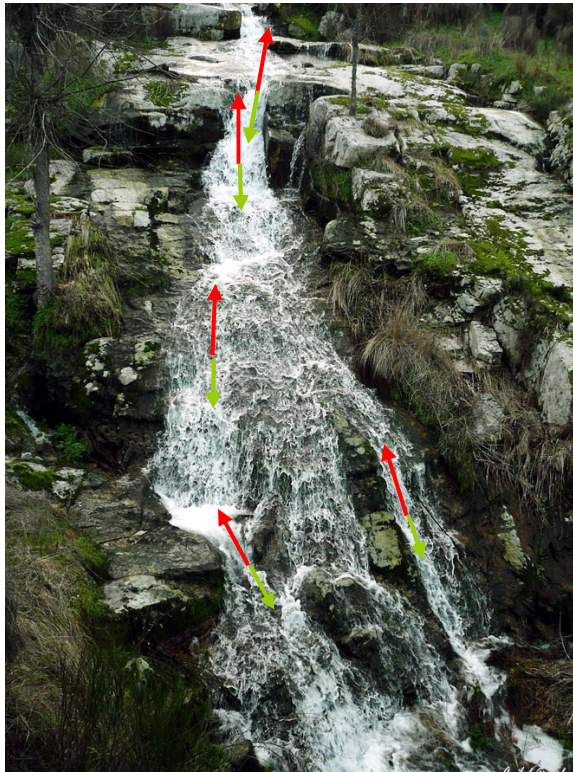
Método da Descida pelo Gradiente

$$\underline{w} \rightarrow \underline{\nabla} \rightarrow \underline{w} \rightarrow \underline{w} - \alpha \underline{\nabla}$$

$$\Delta \underline{w} = - \alpha \underline{\nabla}$$



Gradiente Descendente



Algoritmo

Até que o critério de parada seja satisfeito

Para cada variável w_i

Calcular $\frac{\partial F}{\partial w_i}$

Fazer $w_{i\text{ novo}} = w_{i\text{ antigo}} - \alpha \frac{\partial F}{\partial w_i}$

Fim (natural) do processo

$$\underline{w} = \underline{w}_0$$

$$E(\Delta \underline{w}) = \underline{0} \quad \Rightarrow \quad \nabla(\underline{w}_0) = \underline{0}$$

Passo de Treinamento α

pequeno (decisão local)

valores típicos para RNs $.01 < \alpha < .1$

Em nosso caso, a função à minimizar é

Dispersão total intra classe para todas as classes:

$$F_{in} = \sum_{\forall C_j} F_j$$

$$F_j = \sum_{\forall \vec{x} \in C_j} \left| \vec{x} - \vec{m}_j \right|^2 = \sum_{\forall \vec{x} \in C_j} \sum_{l=1}^n \left(x_l - m_{jl} \right)^2$$

$$\frac{\partial F_{in}}{\partial m_{jl}} = -2 \sum_{\forall \vec{x} \in C_j} \left(x_l - m_{jl} \right)$$

Solução analítica

$$\left[\nabla F_{in} = \frac{\partial F_{in}}{\partial m_{jl}} \right] = -2 \sum_{\forall \vec{x} \in C_j} (x_l - m_{jl}) \Big] = 0]$$

$$\sum_{\forall \vec{x} \in C_j} (x_l - m_{jl}) = 0 \quad \sum_{\forall \vec{x} \in C_j} x_l - n_j m_{jl} = 0 \quad m_{jl} = \frac{1}{n_j} \sum_{\forall \vec{x} \in C_j} x_l = \frac{E}{n_j} x_l$$

$$\vec{m}_j = \frac{1}{n_j} \sum_{\forall \vec{x} \in C_j} \vec{x} = \frac{E}{n_j} \vec{x}$$

Solução Numérica Recursiva

$$\Delta m_{ij} = -\alpha \frac{\partial F_{in}}{\partial m_{jl}} = 2\alpha \sum_{\forall \vec{x} \in C_j} (x_l - m_{jl})$$

$$\Delta \vec{m}_j = 2\alpha \sum_{\forall \vec{x} \in C_j} (\vec{x} - \vec{m}_j)$$

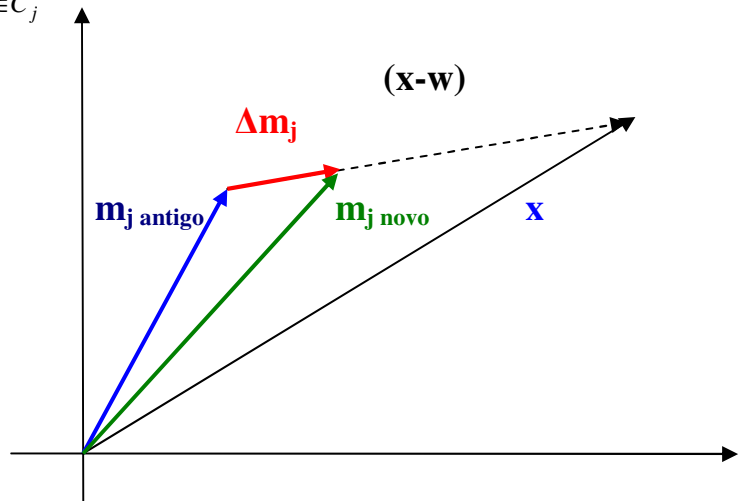
Este é um processo por batelada.

Alternativa: Calcular o acréscimo devido à cada entrada e aplicar imediatamente (processo por “regra delta”). Plástico e reduz a quantidade de memória necessária.

Processo por “regra delta”

$$\Delta m_{ij} = -\alpha \frac{\partial F_{in}}{\partial m_{jl}} = 2\alpha (x_l - m_{jl}) \Big|_{\forall \vec{x} \in C_j}$$

$$\Delta \vec{m}_j = 2\alpha (\vec{x} - \vec{m}_j) \Big|_{\forall \vec{x} \in C_j}$$



Fim do processo recursivo:

$$E(\Delta \vec{m}_j) = 2\alpha E(\vec{x} - \vec{m}_j) \Big|_{\forall \vec{x} \in C_j} = \vec{0}$$

$$\vec{m}_j = \frac{1}{n_j} \sum_{\forall \vec{x} \in C_j} \vec{x} = E_{\forall \vec{x} \in C_j} \vec{x}$$