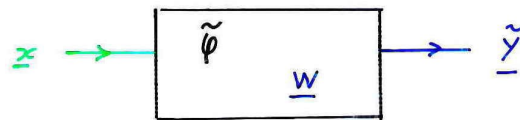


12 -Outros critérios de erro

Composição de F como erro



$$F(\underline{y}, \underline{\tilde{y}}) \quad \underline{\tilde{y}}(\underline{w}) \quad \frac{\partial F}{\partial w_i} = \frac{\partial F}{\partial \tilde{y}} \frac{\partial \tilde{y}}{\partial w_i}$$

onde $\frac{\partial F}{\partial \tilde{y}}$ depende da função erro utilizada e
 $\frac{\partial \tilde{y}}{\partial w_i}$ depende das equações da rede

no caso de saídas múltiplas $\frac{\partial F}{\partial w_i} = \sum_{l=1}^m \frac{\partial F}{\partial \tilde{y}_l} \frac{\partial \tilde{y}_l}{\partial w_i}$

Caso geral: Erro médio quadrático ponderado

$$F = E_k \left\{ \sum_l a(k, l) \varepsilon_l^2 \right\}$$

onde $a(l)$ peso variável por saída l

$a(k)$ peso variável por par k

como $\varepsilon_l = y_l - \tilde{y}_l$

$$\frac{\partial F}{\partial \tilde{y}} = -2 E_k \left\{ \sum_l a(k, l) \varepsilon_l \right\}$$

No treinamento tipo regra delta o acréscimo na sinapse w_{ij} sera dado por

$$\Delta w_{ij} = 2\alpha v_j \delta_i \quad \text{onde} \quad \delta_i = \sum_{l=1}^m a(k,l) \varepsilon_l g_{li}$$

onde

v_j é o sinal na entrada da sinapse na fase de propagação do sinal para a frente e

$g_{li} = \frac{\partial \tilde{y}_l}{\partial u_i}$ é o ganho linearizado da saída da sinapse w_{ij} à saída l da rede

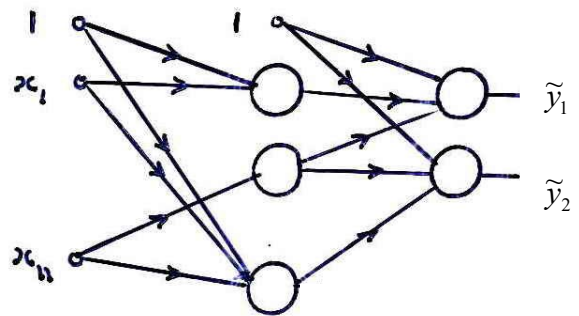
Ex 1 : Ponderando diferentemente saídas diferentes

A rede tem duas saídas, mas

o erro em y_1 é 3 vezes mais importante que o erro em y_2

$$a(k,1) = 3$$

$$a(k,2) = 1$$



$$F = E_k (3\varepsilon_1^2 + \varepsilon_2^2)$$

$$\Delta w_{ij} = 2\alpha v_j \delta_i \quad \text{onde} \quad \delta_i = 3\varepsilon_1 g_{1i} + \varepsilon_2 g_{2i}$$

Ex 2 : Ponderando diferentemente em função da saída desejadas

Falsos negativos x falsos positivos com pesos diferentes

$$F = \frac{1}{4P} \sum_{\forall p} a(y) \varepsilon^2 = \frac{1}{4P} \sum_{\forall p} [(M-1)y + (M+1)](y - \tilde{y})^2 =$$

$$= \frac{1}{4P} \left[\sum_{\substack{\forall p|y=1 \\ \forall \text{ pares positivos}}} 2M(y - \tilde{y})^2 + \sum_{\substack{\forall p|y=-1 \\ \forall \text{ pares negativos}}} (y - \tilde{y})^2 \right]$$

Tipo	y	$\tilde{y}$	contribuição para F
verdadeiros positivos	1	1	0
verdadeiros negativos	-1	-1	0
falsos positivos	-1	1	1
falsos negativos	1	-1	2M

casos:
doenças humanas,
tuberculose bovina,
etc.

Ex 3: Erro relativo médio quadrático

Erro relativo médio quadrático nas variáveis originais, não escaladas

$$F(\underline{w}) = E \left\{ \sum_l \left[\frac{Y_l - \tilde{Y}_l}{Y_l} \right]^2 \right\}$$

$$y_l = \frac{1}{\sigma_{Y_l}} (Y_l - \mu_{Y_l})$$

$$Y_l = \sigma_{Y_l} y_l + \mu_{Y_l}$$

$$\tilde{Y}_l = \sigma_{Y_l} \tilde{y}_l + \mu_{Y_l}$$

$$F(\underline{w}) = E \left\{ \sum_l \left[\frac{Y_l - \tilde{Y}_l}{Y_l} \right]^2 \right\} = E \left\{ \sum_l \frac{1}{\left(y_l + \frac{\mu_{Y_l}}{\sigma_{Y_l}} \right)^2} (y_l - \tilde{y}_l)^2 \right\}$$

$$a(k, l) = \frac{1}{\left(y_l + \frac{\mu_{Y_l}}{\sigma_{Y_l}} \right)^2}$$

No treinamento tipo BP o acréscimo na sinapse w_{ij} é dado por

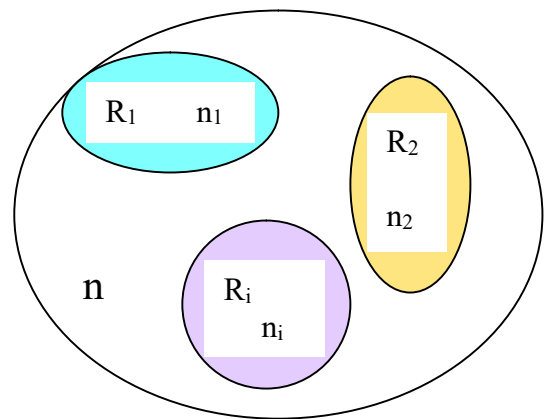
$$\Delta w_{ij} = 2\alpha v_j \delta_i \quad \text{onde} \quad \delta_i = \sum_{l=1}^m \frac{1}{\left(y_l + \frac{\mu_{Y_l}}{\sigma_{Y_l}}\right)^2} \varepsilon_l g_{li}$$

Ex 4: Correção do efeito da população local

Cada par pertence a uma região R_i

Cada região R_i tem uma população n_i

A população total é $n = \sum_{\forall i} n_i$



$$F = \sum_{\forall \text{ região } R_i} \sum_{\forall \text{ par} \in R_i} f(y, \tilde{y})$$

peso: n_i pares

$$F = \sum_{\forall \text{ região } R_i} \frac{1}{n_i} \sum_{\forall \text{ par} \in R_i} f(y, \tilde{y})$$

peso: 1 par, independente de n_i

$$a(i, l) = a(i) = 1/n_i$$

n_i – população da região onde o par pertence.

12.2 Erros não quadráticos

12.2.1 Erros em potências maiores

$$F(\underline{w}) = E \left\{ \sum_l (y_l - \tilde{y}_l)^{2n+2} \right\} \quad n = 1, 2, \dots$$

prioriza a redução dos maiores erros, mas $F(w)$ é mais abrupta.

E o minimante ? Caso uma saída

$$F = E(y - \tilde{y})^{2n+2}$$

$$\frac{\partial F}{\partial \tilde{y}} = (2n+2) E(y - \tilde{y})^{2n+1} = 0$$

Usando propriedades da expansão binomial é possível provar que a equação acima é satisfeita (para $2n+1$ ímpar) se $E(y) = E(\tilde{y})$

12.2.2 Erro absoluto

$$F(\underline{w}) = E_k \left\{ | \vec{y}_k - \tilde{\vec{y}}_k | \right\} = E_k \left\{ \sum_l | y_{kl} - \tilde{y}_{kl} | \right\} = E_k \left\{ \sum_l | \varepsilon_{kl} | \right\}$$

não prioriza a redução dos maiores erros como o emq

- não derivável na origem, necessita procedimentos especiais

$$F(\underline{w}) = E_k \left\{ \left| \vec{y}_k - \tilde{\vec{y}}_k \right| \right\} = E_k \left\{ \sum_l \left| y_{kl} - \tilde{y}_{kl} \right| \right\} = E_k \left\{ \sum_l \left| \varepsilon_{kl} \right| \right\}$$

$$\left| \varepsilon_{kl} \right| = \begin{cases} \varepsilon_{kl} = y_{kl} - \tilde{y}_{kl} & se \quad \varepsilon_{kl} > 0 \\ -\varepsilon_{kl} = -(y_{kl} - \tilde{y}_{kl}) & se \quad \varepsilon_{kl} < 0 \end{cases}$$

$$\frac{\partial}{\partial w} \left| \varepsilon_{kl} \right| = \begin{cases} -\frac{\partial \tilde{y}_{kl}}{\partial w} & se \quad \varepsilon_{kl} > 0 \\ \frac{\partial \tilde{y}_l}{\partial w} & se \quad \varepsilon_{kl} < 0 \end{cases}$$

$$\frac{\partial}{\partial w} \left| \varepsilon_{kl} \right| = -a(k, l) \frac{\partial \tilde{y}_{kl}}{\partial w} = -sign(\varepsilon_{lk}) \frac{\partial \tilde{y}_{kl}}{\partial w}$$

12.2.3 Erro MAPE – Mean Absolute Percentual Error

$$F(\underline{w}) = E_k \left\{ \frac{\left| \vec{Y} - \tilde{\vec{Y}} \right|}{\left| \vec{Y} \right|} \right\}$$

variáveis não escaladas

não prioriza a redução dos maiores erros percentuais, como o epmq.

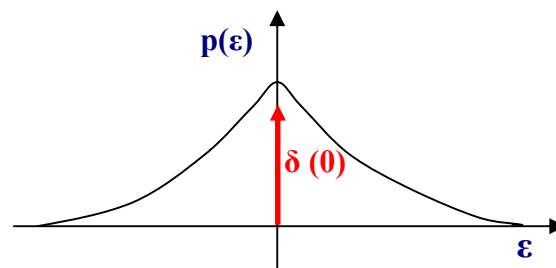
- não derivável na origem, necessita procedimentos especiais

12.3 Entropia

O objetivo na minimização do erro ε é

obter o histograma de $p(\varepsilon)$ com média nula (fácil) e o mais estreito possível

ideal: $p(\varepsilon) = \delta(0)$ (delta de Dirac, erro nulo)

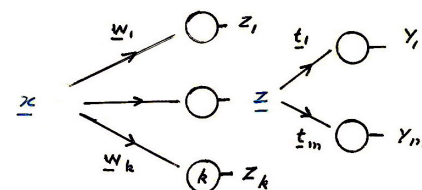


Se a distribuição do erro $p(\varepsilon)$ for Gaussiana, o objetivo é alcançado minimizando o erro médio quadrático (a variância de $p(\varepsilon)$)

Se a distribuição do erro $p(\varepsilon)$ não for Gaussiana, o objetivo é alcançado minimizando a entropia (J.C. Príncipe)

Entropia Relativa (à minimizar)

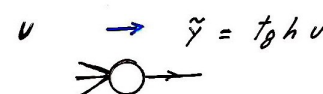
$$F_{ER} = E \sum_{\forall \text{ pares } l} \left[(1 + y_l) \ln \frac{1 + y_l}{1 + \tilde{y}_l} + (1 - y_l) \ln \frac{1 - y_l}{1 - \tilde{y}_l} \right]$$



Por BP regra delta

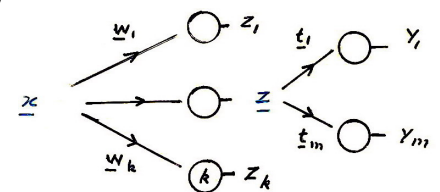
$$\Delta w = -\alpha \frac{\partial f}{\partial w} = -\alpha \frac{\partial f}{\partial \tilde{y}} \frac{\partial \tilde{y}}{\partial u} \frac{\partial u}{\partial w}$$

Camada de saída



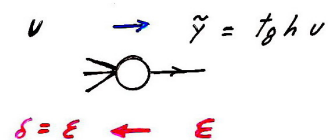
Para a entropia relativa, após alguma álgebra

$$\begin{aligned}\Delta w_{ER} &= -\alpha \frac{\partial}{\partial w} \left[(1 + y_l) \ln \frac{1 + y_l}{1 + \tilde{y}_l} + (1 - y_l) \ln \frac{1 - y_l}{1 - \tilde{y}_l} \right] = \\ &= -\alpha \frac{\partial}{\partial \tilde{y}} \left[(1 + y_l) \ln \frac{1 + y_l}{1 + \tilde{y}_l} + (1 - y_l) \ln \frac{1 - y_l}{1 - \tilde{y}_l} \right] \frac{\partial \tilde{y}}{\partial u} \frac{\partial u}{\partial w} = \\ &= -\alpha \left[-2 \frac{y - \tilde{y}}{1 - \tilde{y}^2} \right] (1 - \tilde{y}^2) \frac{\partial u}{\partial w} = 2\alpha \varepsilon \frac{\partial u}{\partial w}\end{aligned}$$



$$\Delta w_{ER} = 2\alpha \varepsilon \frac{\partial u}{\partial w}$$

Camada de saída

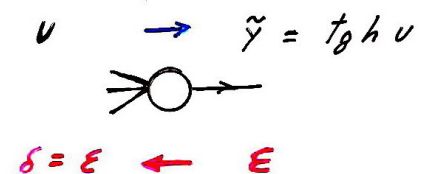


Comparando com o EMQ

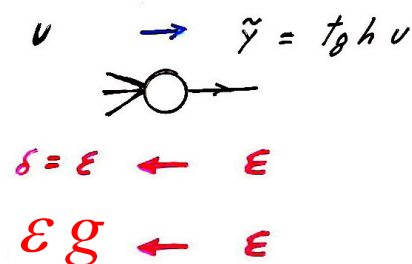
Retropropagação do erro

$$\Delta w_{EMQ} = 2\alpha \varepsilon g \frac{\partial u}{\partial w} \quad \Delta w_{ER} = \hat{\varepsilon} \frac{\partial u}{\partial w}$$

Como $\Delta w_{EMQ} = g \Delta w_{ER}$ o minimant convergência é diferente.



Camada de saída

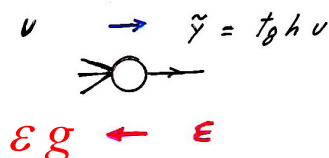


Retropropagação do erro: a única diferença é que para minimizar a ER usa-se

$g=1$ (mas apenas para a camada de saída)

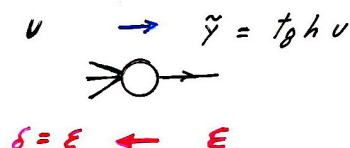
Camada de saída

EMQ



Camada de saída

ER



Retropropagação do erro – ganho dos neurônios

Redes com uma camada

Camadas de saída com neurônios tgh

Ganho dos neurônios da camada de saída $g_i = 1$

Camada de saída com neurônios lineares

Ganho dos neurônios da camada de saída $g_i = \frac{1}{(1 - \tilde{y}_i^2)}$

Redes com duas camadas

Camadas intermediária e de saída com neurônios tgh

Ganho dos neurônios de saída $g_i = 1$

Ganho dos neurônios da camada intermediária $g_i = (1 - z_i^2)$

Camada intermediária neurônios tgh e camada de saída com neurônios lineares

Ganho dos neurônios de saída $g_i = \frac{1}{(1 - \tilde{y}_i^2)}$

Ganho dos neurônios da camada intermediária $g_i = (1 - z_i^2)$

12.4 - Erro lógico (já visto em classificadores)

$$y \in \{-1, +1\} \quad \tilde{y} = \tanh(u) \in (-1, +1)$$

$$\tilde{y}_{\text{lógico}} = \text{sign}(\tilde{y})$$

Erro médio quadrático

$$\frac{1}{N} \sum_i (y_i - \tilde{y}_i)^2$$

Erro de classificação

$$\frac{1}{4N} \sum_{i=1}^N (y_i - \tilde{y}_{i \text{ lógico}})^2$$

% de classificações erradas em y_i