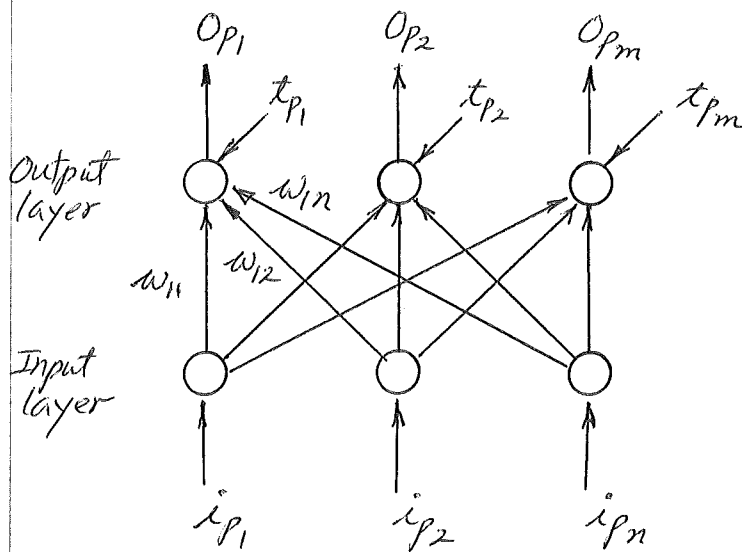


Delta Rule

p : pattern

i_p : input vector for pattern p

O_p : output vector for pattern p

t_p : target vector for pattern p

w_{ij} : weight from j th input to i th output component.

For a pattern p , the error is

$$E_p = \frac{1}{2} \sum_j (t_{pj} - O_{pj})^2, \quad E = \sum_p E_p$$

The gradient is, by chain rule,

$$\frac{\partial E_p}{\partial w_{ji}} = \frac{\partial E_p}{\partial O_{pj}} \frac{\partial O_{pj}}{\partial w_{ji}}$$

$$\frac{\partial E_p}{\partial O_{pj}} = -(t_{pj} - O_{pj}) \triangleq -\delta_{pj}$$

Assuming linear units,

$$O_{pj} = \sum_i w_{ji} i_{pi} \Rightarrow \frac{\partial O_{pj}}{\partial w_{ji}} = i_{pi}$$

Therefore, the negative gradient is

$$-\frac{\partial E_p}{\partial w_{ji}} = \delta_{pj} i_{pi}$$

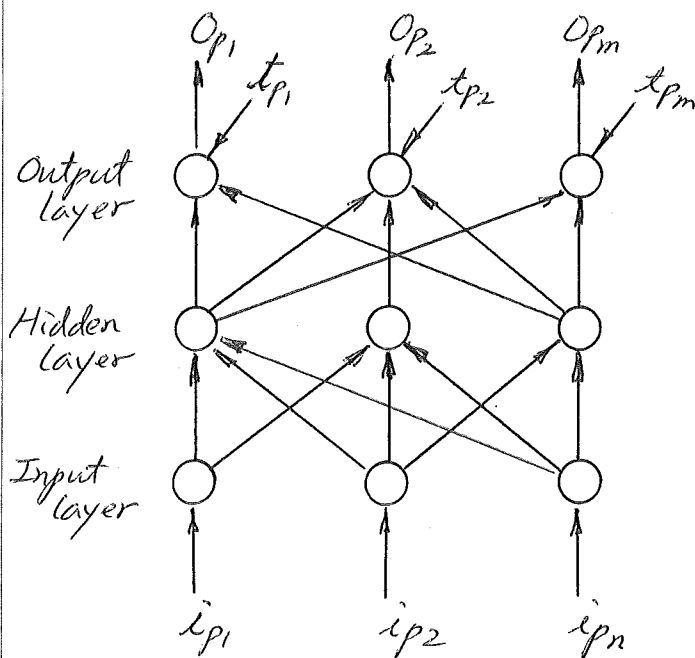
$\therefore$ The weight w_{ji} can be updated by

$$\Delta_p w_{ji} = \eta \delta_{pj} i_{pi}, \quad \eta = \text{learning rate}$$

$$\text{Since } \frac{\partial E}{\partial w_{ji}} = \sum_p \frac{\partial E_p}{\partial w_{ji}},$$

the Delta rule implements a gradient descent on E .

Generalized Delta Rule



Now we add hidden layers (can be multiple).

Assume semilinear activation function:

$$Op_j = f_j(\text{net}_{pj}) \quad \dots (1)$$

$$\text{net}_{pj} = \sum_i w_{ji} Op_i \quad \dots (2)$$

f_j : differentiable, nondecreasing

net_{pj} : weighted sum of inputs to j th neuron for pattern p .

Note: inputs to j th neuron are outputs of i th neuron in the previous layer.

Define the error function:

$$E_p = \frac{1}{2} (tp_j - Op_j)^2, \quad E = \sum_p E_p$$

Then the gradient is

$$\frac{\partial E_p}{\partial w_{ji}} = \underbrace{\frac{\partial E_p}{\partial \text{net}_{pj}}}_{Op_i \text{ from (2)}} \underbrace{\frac{\partial \text{net}_{pj}}{\partial w_{ji}}}_{\Rightarrow -\frac{\partial E_p}{\partial w_{ji}} = \delta_{pj} Op_i \dots (3)}$$

$$\begin{aligned} \text{Define } \delta_{pj} &\triangleq -\frac{\partial E_p}{\partial \text{net}_{pj}} \\ &= -\frac{\partial E_p}{\partial Op_j} \underbrace{\frac{\partial Op_j}{\partial \text{net}_{pj}}}_{f'_j(\text{net}_{pj}) \text{ from (1)}} \quad \dots (4) \end{aligned}$$

Now we have two different cases to compute the first term: $\frac{\partial E_p}{\partial Op_j}$

Case I. Output layer

$$\frac{\partial E_p}{\partial O_j} = -(t_j - O_j)$$

Thus, from (4),

$$\delta_{p_j} = (t_j - O_j) f'_j(\text{net}_{p_j}) \quad (5)$$

Case II. Hidden layers

Note that when the output O_j of hidden layer changes the net inputs net_{pk} to the output layer changes.

$$\text{net}_{pk} = \sum_i w_{ki} O_i$$

Thus,

$$\begin{aligned} \frac{\partial E_p}{\partial O_j} &= \sum_k \frac{\partial E_p}{\partial \text{net}_{pk}} \underbrace{\frac{\partial \text{net}_{pk}}{\partial O_j}}_{w_{kj}} \\ &= \sum_k \underbrace{\frac{\partial E_p}{\partial \text{net}_{pk}}}_{-\delta_{pk}} w_{kj} \end{aligned}$$

$-\delta_{pk}$: already calculated at the output layer.

$$= - \sum_k \delta_{pk} w_{kj}$$

Thus, from (4),

$$\delta_{p_j} = f'_j(\text{net}_{p_j}) \sum_k \delta_{pk} w_{kj} \quad \text{error is propagated backward} \quad (6)$$

The weight adjustment is, therefore, in the direction of the negative gradient:

$$\Delta_p w_{ji} = \eta \delta_{p_j} O_i \quad \dots \dots \dots (7)$$

NOTE:

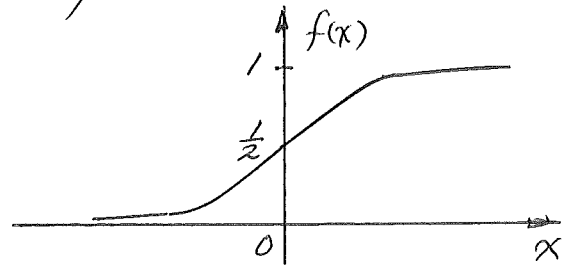
δ_{p_j} is propagated backward (to neuron j from the next layer).
 O_i is propagated forward (from neuron i in the previous layer)

Note that the errors in the output layer (5) are propagated backward through output weights (weights between output layer and hidden layer) (6), and multiplied by inputs propagated forward (7) to compute the weight adjustment.

Derivative of the activation function, $f'_j(\text{net}p_j)$, can be computed directly without differentiating numerically.

For example, for a sigmoidal function

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



it can be shown that

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

Activation function ($\tau = 0$)

Therefore, from (1),

$$f'_j(\text{net}p_j) = O_{pj} [1 - O_{pj}]$$

In general, the momentum is added to avoid local minima:

$$\Delta_p w_{ji}^{\text{new}} = \eta \delta p_j O_{pi} + \alpha \Delta_p w_{ji}^{\text{prev}}$$

where the second term on the right hand side is the momentum term. This term adds a portion of the most recent weight change when computing the new weight change. The momentum term is supposed to give the neuron momentum in weight space, enabling it to pass through local minima.