

Module II

Perception networks – Learning rule – Training and testing algorithm, Adaptive Linear Neuron, Back propagation Network – Architecture, Training algorithm.

Perception networks

Theory

Perception networks come under single-layer feed-forward networks and are also called **simple perceptrons**.

The key points to be noted in a perceptron network are:

1. The perceptron network consists of three units, namely, sensory unit (input unit), associator unit (hidden unit), response unit (output unit).
2. The sensory units are connected to associator units with fixed weights having values 1, 0 or -1, which are assigned at random.
3. The binary activation function is used in sensory unit and associator unit.
4. The response unit has an activation of 1, 0 or -1. The binary step with fixed threshold θ is used as activation for associator. The output signals that are sent from the associator unit to the response unit are only binary.
5. The output of the perceptron network is given by

$$y = f(y_{in})$$

where $f(y_{in})$ is activation function and is defined as

$$f(y_{in}) = \begin{cases} 1 & \text{if } y_{in} > \theta \\ 0 & \text{if } -\theta \leq y_{in} \leq \theta \\ -1 & \text{if } y_{in} < -\theta \end{cases}$$

6. The perceptron learning rule is used in the weight updation between the associator unit and the response unit. For each training input, the net will calculate the response and it will determine whether or not an error has occurred.

7. The error calculation is based on the comparison of values of targets with those of the calculated outputs.

8. The weights on the connections from the units that send the nonzero signal will get adjusted suitably.

9. The weights will be adjusted on the basis of the learning rule if an error has occurred for a particular training pattern, i.e.,

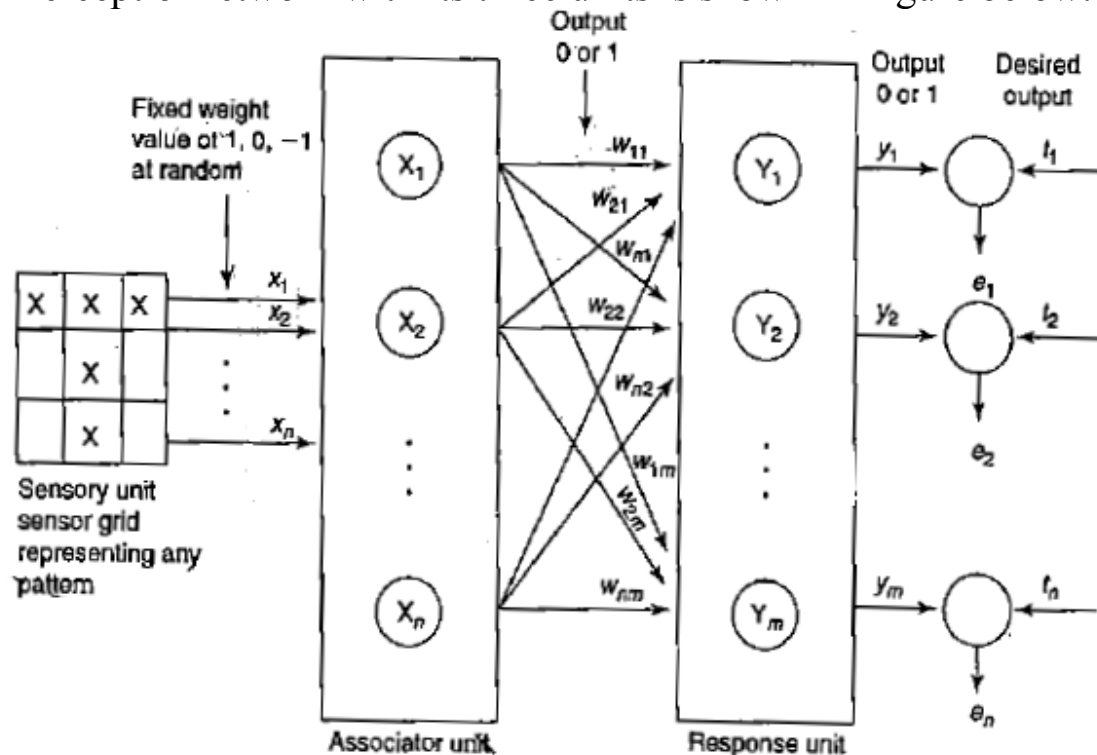
$$w_i(\text{new}) = w_i(\text{old}) + \alpha \delta_i$$

$$b(\text{new}) = b(\text{old}) + \alpha r$$

If no error occurs, there is no weight updation and hence the training process may be stopped. In the above equations, the target value "t" is +1 or -1 and α is the learning rate.

Original Perception network

A Perception network with its three units is shown in Figure below:



Sensory unit

A sensory unit can be a two-dimensional matrix of 400 photo detectors upon which a lighted picture with geometric black and white pattern impinges. These detectors provide a binary (0 or 1) electrical signal if the input signal is found to exceed a certain value of threshold. Also, these detectors are connected randomly with the associator unit.

Associator unit

The associator unit is found to consist of a set of subcircuits called **feature predicates**. The feature predicates are hard-wired to detect the specific feature of a pattern and are equivalent to the *feature detectors*. For a particular feature, each predicate is examined with a few or all of the responses of the sensory unit. It can be found that the results from the predicate units are also binary (0 or 1).

Response unit

The last unit, i.e. response unit, contains the pattern recognizers or perceptrons. The weights present in the input layers are all fixed, while the weights on the response unit are trainable.

Learning rule

In case of the perceptron learning rule, the learning signal is the difference between the desired and actual response of a neuron. The perceptron learning rule is explained as follows:

Consider a finite " n " number of input training vectors, with their associated target values $x(n)$ and $t(n)$, where " n " ranges from 1 to N . The target is either +1 or -1. The output " y " is obtained on the basis of the net input calculated and activation function being applied over the net input.

$$y = f(y_{in}) = \begin{cases} 1 & \text{if } y_{in} > \theta \\ 0 & \text{if } -\theta \leq y_{in} \leq \theta \\ -1 & \text{if } y_{in} < -\theta \end{cases}$$

The weight updation in case of perceptron learning is as shown.

If $y \neq t$ then

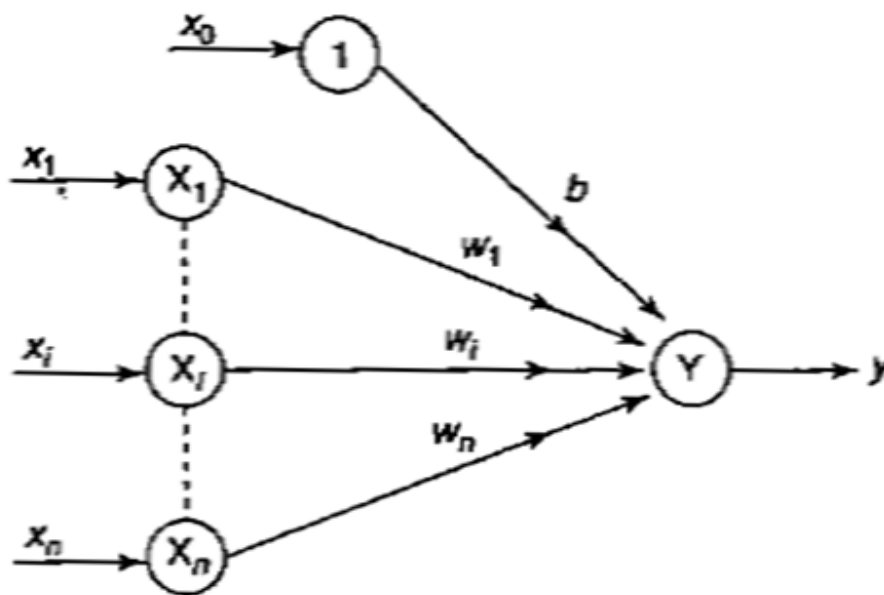
$$w(\text{new}) = w(\text{old}) + \alpha tx \quad (\alpha - \text{learning rate})$$

else

$$w(\text{new}) = w(\text{old})$$

Architecture

A simple perceptron network architecture is shown in Figure below:

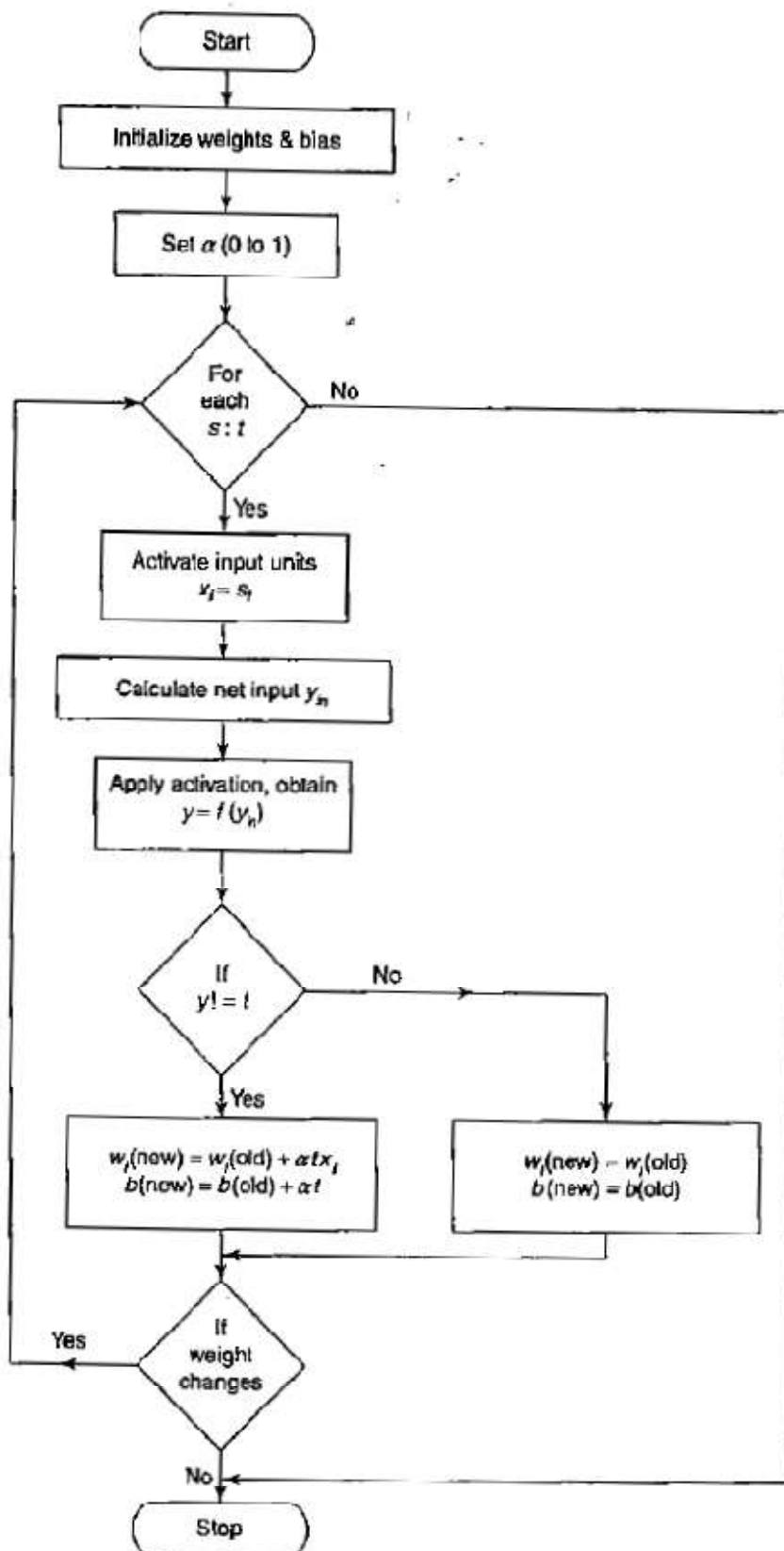


In Figure, there are n input neurons, 1 output neuron and a bias. The input-layer and output layer neurons are connected through a directed communication link, which is associated with weights.

The goal of the perceptron net is to classify the input pattern as a member or not a member to a particular class.

Flowchart for Training Process

The flowchart for the perceptron network training is shown in Figure.



The flowchart depicted here presents the flow of the training process. As depicted in the flowchart, first the basic initialization required for the training process is performed.

The entire loop of the training process continues until the training input pair is presented to the network. The training (weight updation) is done on the basis of the comparison between the calculated and desired output. The loop is terminated if there is no change in weight.

Training and testing algorithm

Perceptron Training Algorithm for Single Output Classes

The perceptron algorithm can be used for either binary or bipolar input vectors, having bipolar targets, threshold being fixed and variable bias.

In the algorithm below, initially the inputs are assigned. Then the net input is calculated. The output of the network is obtained by applying the activation function over the calculated net input.

On performing comparison over the calculated and the desired output, the weight updation process is carried out. The entire network is trained based on the mentioned stopping criterion.

The algorithm of a perceptron network is as follows:

Step 0: Initialize the weights and the bias. Also initialize the learning rate α ($0 < \alpha \leq 1$). For simplicity α is set to 1.

Step 1: Perform Steps 2-6 until the final stopping condition is false.

Step 2: Perform Steps 3-5 for each training pair indicated by $s:t$.

Step 3: The input layer containing input units is applied with identity activation functions:

$$x_i = f_i$$

Step 4: Calculate the output of the network. To do so, first obtain the net input:

$$y_{in} = b + \sum_{i=1}^n x_i w_i$$

where "n" is the number of input neurons in the input layer. Then apply activations over the netinput calculated to obtain the output:

$$y = f(y_{in}) = \begin{cases} 1 & \text{if } y_{in} > \theta \\ 0 & \text{if } -\theta \leq y_{in} \leq \theta \\ -1 & \text{if } y_{in} < -\theta \end{cases}$$

Step 5:Weight and bias adjustment: Compare the value of the actual (calculated) output and desired(target) output.

If $y \neq t$, then

$$w_i(\text{new}) = w_i(\text{old}) + \alpha x_i$$

$$b(\text{new}) = b(\text{old}) + \alpha t$$

else, we have

$$w_i(\text{new}) = w_i(\text{old})$$

$$b(\text{new}) = b(\text{old})$$

Step 6: Train the network until there is no weight change. This is the stopping condition for the network.If this condition is not met, then start again from Step 2.

Perceptron Training Algorithm for Multiple Output Classes

For multiple output classes, the perceptron training algorithm is as follows:

Step 0:Initialize the weights, biases and learning rate suitably.

Step 1: Check for stopping condition; if it is false, perform Steps 2-6.

Step 2: Perform Steps 3--5 for each bipolar or binary training vector pair $s:t$.

Step 3: Set activation (identity) of each input unit $i= 1$ to n :

$$X_i = S_i$$

Step 4: Calculate output response of each output unit $j=1$ to m ; First the net input is calculated as:

$$y_{inj} = b_j + \sum_{i=1}^n x_i w_{ij}$$

Then activations are applied over the net input to calculate the output response:

$$y_j = f(y_{inj}) = \begin{cases} 1 & \text{if } y_{inj} > \theta \\ 0 & \text{if } -\theta \leq y_{inj} \leq \theta \\ -1 & \text{if } y_{inj} < -\theta \end{cases}$$

Step 5: Make adjustment in weights and bias for $j=1$ to m and $i=1$ to n .

If $t_j \neq y_j$, then

$$w_{ij}(\text{new}) = w_{ij}(\text{old}) + \alpha t_j x_i$$

$$b_j(\text{new}) = b_j(\text{old}) + \alpha t_j$$

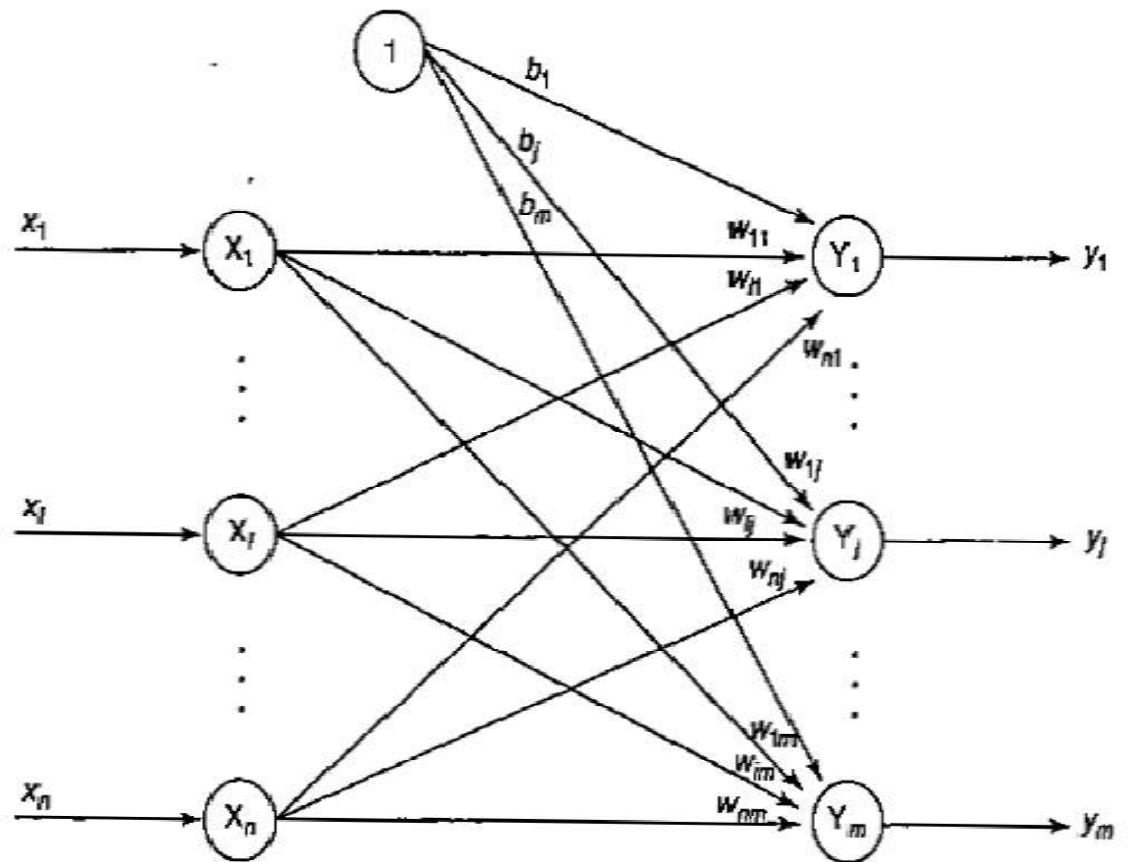
else, we have

$$w_{ij}(\text{new}) = w_{ij}(\text{old})$$

$$b_j(\text{new}) = b_j(\text{old})$$

Step 6: Test for the stopping condition, i.e., if there is no change in weights then stop the training process, else start again from Step 2.

The above algorithm is suited for the architecture shown in Figure below.



Perceptron Network Testing Algorithm

The testing algorithm is as follows:

Step 0: The initial weights to be used here are taken from the training algorithms.

Step 1: For each input vector X to be classified, perform Steps 2-3.

Step 2: Set activations of the input unit.

Step 3: Obtain the response of output unit.

$$y_{in} = \sum_{i=1}^n x_i w_i$$

$$y = f(y_{in}) = \begin{cases} 1 & \text{if } y_{in} > \theta \\ 0 & \text{if } -\theta \leq y_{in} \leq \theta \\ -1 & \text{if } y_{in} < -\theta \end{cases}$$

Adaptive Linear Neuron (Adaline)

Theory

The units with linear activation function are called linear units. A network with a single linear unit is called an **Adaline (adaptive linear neuron)**. That is, in an Adaline, the input-output relationship is linear.

Adaline uses bipolar activation for its input signals and its target output. The weights between the input and the output are adjustable. The bias in Adaline acts like an adjustable weight, whose connection is from a unit with activations being always 1. Adaline is a net which has only one output unit. The Adaline network may be trained using delta rule.

The delta rule may also be called as least mean square (LMS) rule or Widrow-Hoff rule. This learning rule is found to minimize the mean-squared error between the activation and the target value.

Delta Rule for Single Output Unit

The Widrow-Hoff rule is very similar to perceptron learning rule. The delta rule updates the weights between the connections so as to minimize the difference between the net input to the output unit and the target value. The major aim is to minimize the error over all training patterns. This is done by reducing the error for each pattern, one at a time.

The delta rule for adjusting the weight of i th pattern $\{i = 1 \text{ to } n\}$ is

$$\Delta w_i = \alpha(t - y_{in})x_i$$

where

Δw_i is the weight change

α the learning rate

x the vector of activation of input unit

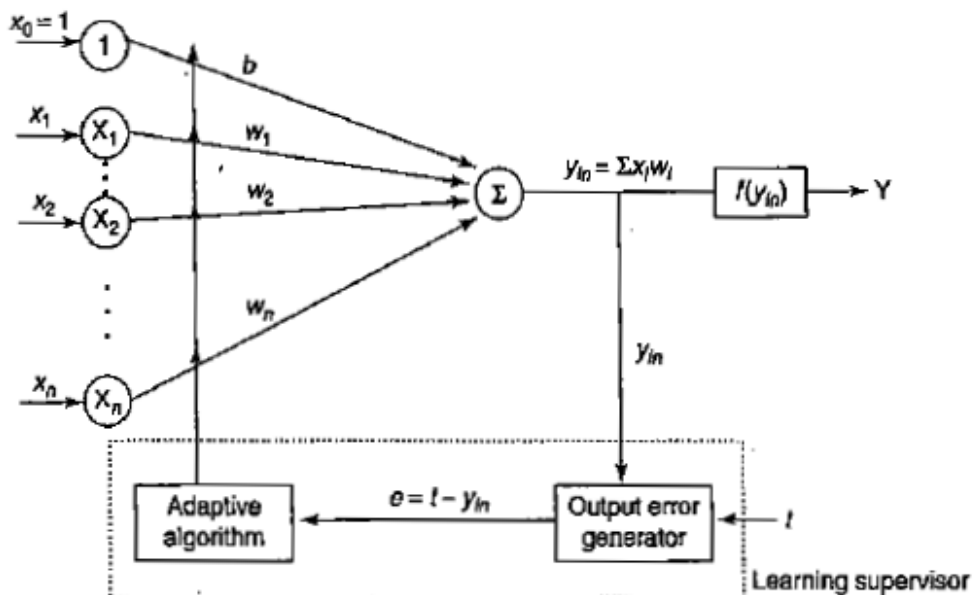
y_{in} the net input to output unit

The delta rule in case of several output units for adjusting the weight from i th input unit to the j th output unit (for each pattern) is

$$\Delta w_{ij} = \alpha(t_j - y_{mj})x_i$$

Architecture

Adaline is a single-unit neuron, which receives input from several units and also from one unit called bias. An Adaline model is shown in Figure below:



The basic Adaline model consists of trainable weights.

Inputs are either of the two values (+ 1 or -1) and the weights have signs (positive or negative).

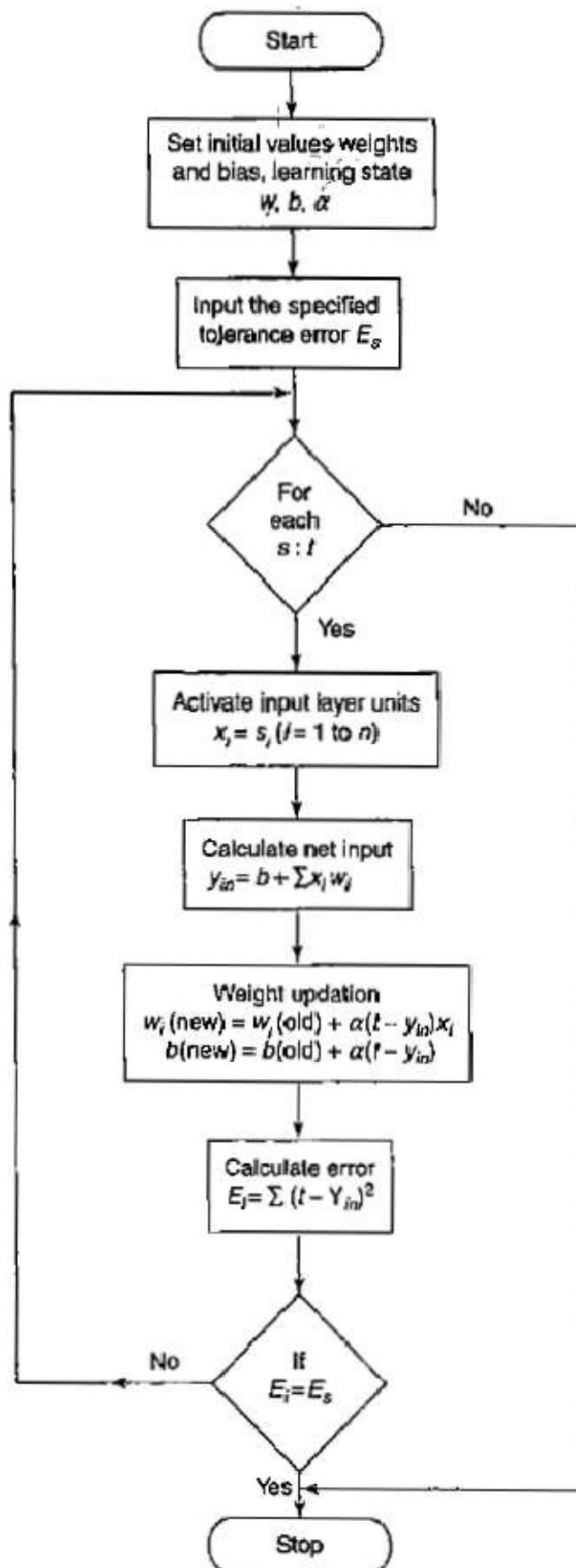
Initially, random weights are assigned.

The net input calculated is applied to a quantizer transfer function that restores the output to +1 or -1. The Adaline model compares the actual output with the target output and on the basis of the training algorithm, the weights are adjusted.

Flowchart for Training Process

The flowchart for the training process is shown in Figure below: This gives a pictorial representation of the network training.

The conditions necessary for weight adjustments have to be checked carefully. The weights and other required parameters are initialized. Then the net input is calculated, output is obtained and compared with the desired output for calculation of error. On the basis of the error Factor, weights are adjusted.



Training Algorithm

The Adaline network training algorithm is as follows:

Step 0: Weights and bias are set to some random values but not zero. Set the learning rate parameter α .

Step 1: Perform Steps 2-6 when stopping condition is false.

Step 2: Perform Steps 3-5 for each bipolar training pair $s: t$.

Step 3: Set activations for input units $i= 1$ to n .

$$x_i = s_i$$

Step 4: Calculate the net input to the output unit.

$$y_{in} = b + \sum_{i=1}^n x_i w_i$$

Step 5: Update the weights and bias for $i= 1$ to n :

$$\begin{aligned} w_i(\text{new}) &= w_i(\text{old}) + \alpha (t - y_{in}) x_i \\ b(\text{new}) &= b(\text{old}) + \alpha (t - y_{in}) \end{aligned}$$

Step 6: If the highest weight change that occurred during training is smaller than a specified tolerance then stop the training process, else continue. This is the rest for stopping condition of a network.

The range of learning rate can be between 0.1 and 1.0.

Testing Algorithm

It is essential to perform the resting of a network that has been trained. When training is completed, the Adaline can be used to classify input patterns.

A step function is used to test the performance of the network. The resting procedure for the Adaline network is as follows:

Step 0: Initialize the weights.

Step 1: Perform Steps 2-4 for each bipolar input vector x .

Step 2: Set the activations of the input units to x .

Step 3: Calculate the net input to the output unit:

$$y_{in} = b + \sum x_i w_i$$

Step 4: Apply the activation function over the net input calculated:

$$y = \begin{cases} 1 & \text{if } y_{in} \geq 0 \\ -1 & \text{if } y_{in} < 0 \end{cases}$$

Back propagation Network

Theory

The backpropagation learning algorithm is one of the most important developments in neural networks. The networks associated with back-propagation learning algorithm are called back propagation networks (BPNs).

For a given set of training input-output pair, this algorithm provides a procedure for changing the weights in a BPN to classify the given input patterns correctly.

The basic concept for this weight update algorithm is simply the gradient-descent method. This is a method where the error is propagated back to the hidden unit.

The aim of the neural network is to train the net to achieve a balance between the net's ability to respond and its ability to give reasonable responses to the input that is similar but not identical to the one that is used in training.

The back-propagation algorithm is different from other networks in respect to the process by which weights are calculated during the learning period of the network.

The general difficulty with the multilayer perceptrons is calculating the weights of the hidden layers in an efficient way that would result in a very small zero output error.

When the hidden layers are increased the network training becomes more complex. To update weights, the error must be calculated. The error, which is the difference between the actual (calculated) and the desired (target) output, is easily measured at the output layer.

The training of the BPN is done in three stages:

- the feed-forward of the input training pattern
- the calculation and back-propagation of the error
- updation of weights.

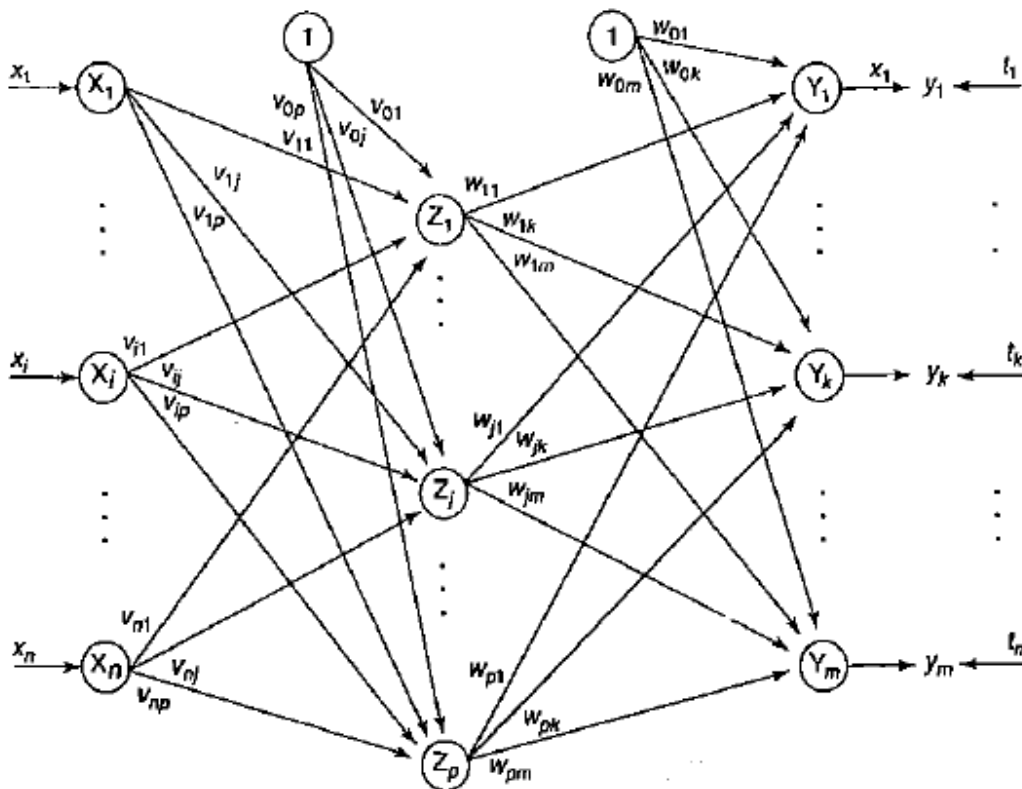
Architecture

A back-propagation neural network is a multilayer, feed forward neural network consisting of an input layer, a hidden layer and an output layer.

The neurons present in the hidden and output layers have biases, which are the connections from the units whose activation is always 1. The bias terms also act as weights.

The Figure below shows the architecture of a BPN, depicting only the direction of information flow for the feed-forward phase. During the back propagation phase of learning, signals are sent in the reverse direction.

The inputs sent to the BPN and the output obtained from the net could be either binary (0, 1) or bipolar (-1, +1). The activation function could be any function which increases monotonically and is also differentiable.



Training algorithm

Flowchart for Training Process

The flowchart for the training process using a BPN is shown in Figure below. The terminologies used in the flowchart and in the training algorithm are as follows:

- x = input training vector $(x_1, \dots, x_i, \dots, x_n)$
- t = target output vector $(t_1, \dots, t_k, \dots, t_m)$
- α = learning rate parameter
- x_i = input unit i . (Since the input layer uses identity activation function, the input and output signals here are same.)
- v_{0j} = bias on j th hidden unit
- w_{0k} = bias on k th output unit
- z_j = hidden unit j . The net input to z_j is

$$z_{inj} = v_{0j} + \sum_{i=1}^n x_i v_{ij}$$

and the output is

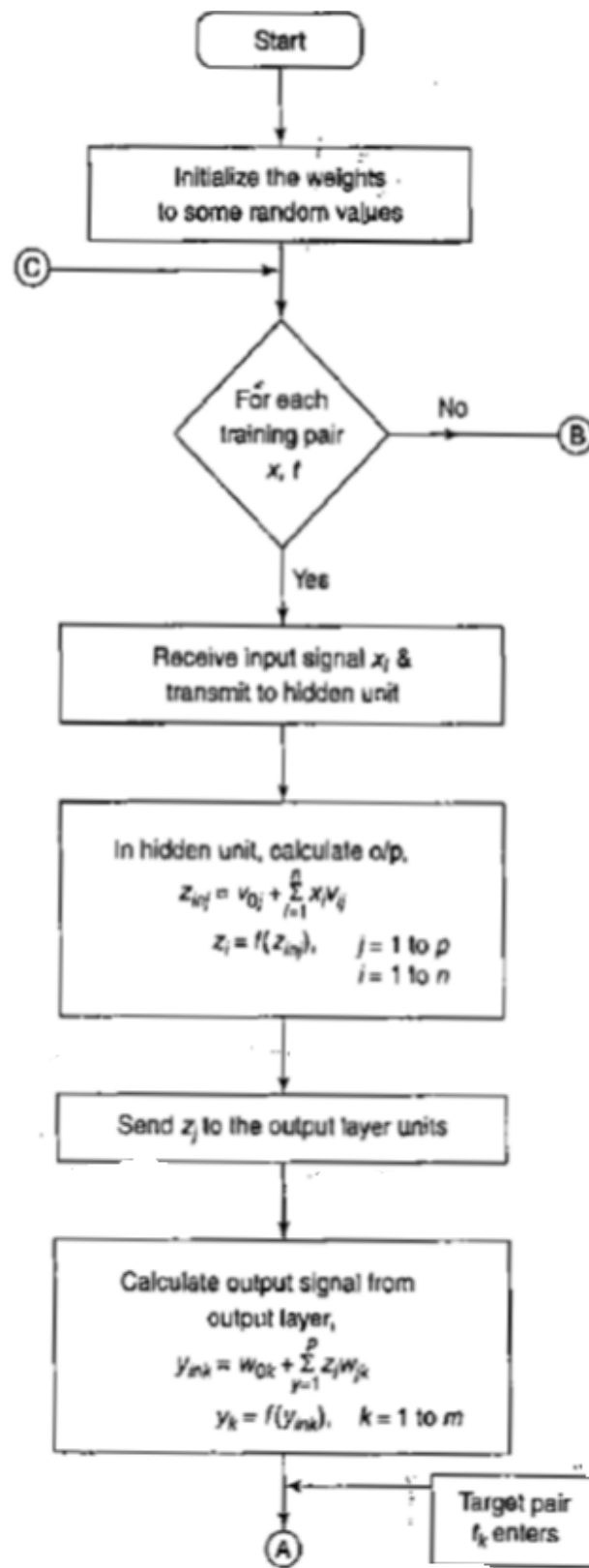
$$z_j = f(z_{inj})$$

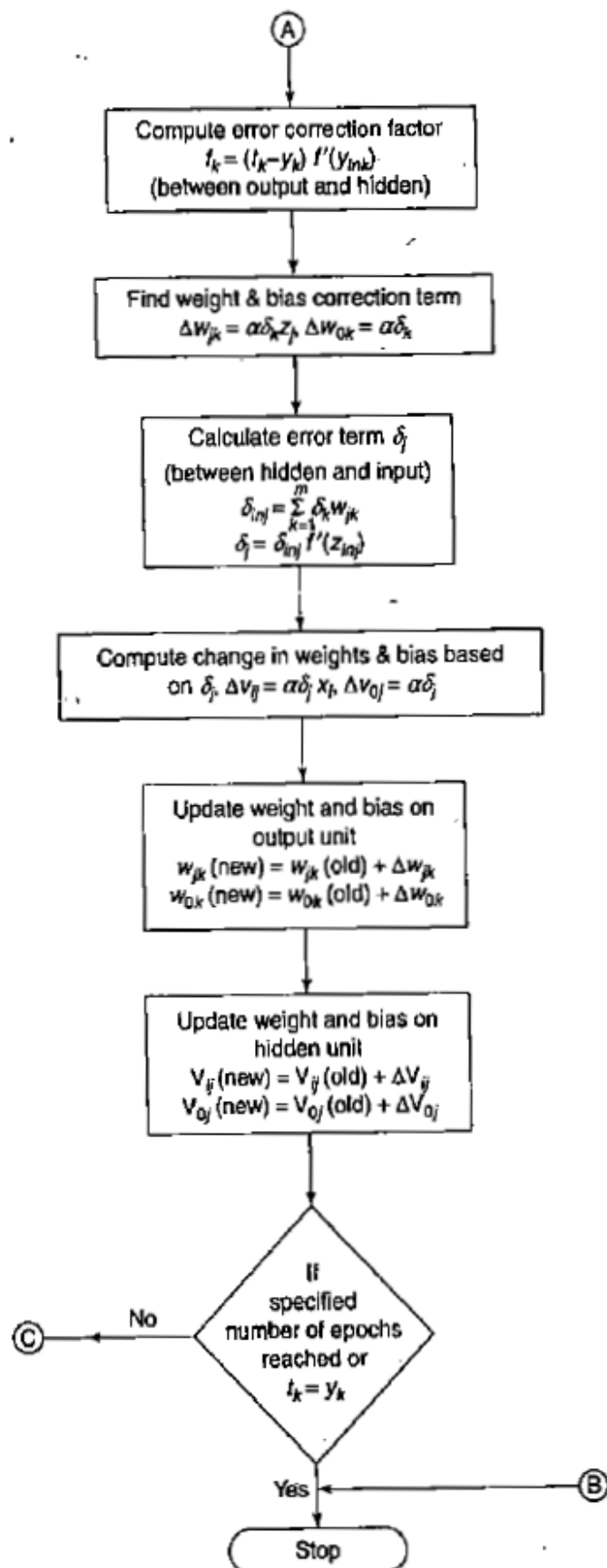
y_k = output unit k . The net input to y_k is

$$y_{ink} = w_{0k} + \sum_{j=1}^p z_j w_{jk}$$

and the output is

$$y_k = f(y_{ink})$$





Training Algorithm

Step 0: Initialize weights and learning rate.

Step 1: Perform Steps 2-9 when stopping condition is false.

Step 2: Perform Steps 3-8 for training pair.

Feed-Forward Phase(Phase: I)

Step 3: Each input unit receives input signal x_i ; and sends it to the hidden unit ($i = 1$ to n).

Step 4: Each hidden unit $z_j(j = 1$ to $p)$ sums its Weighted input signals to calculate net input:

$$z_{inj} = w_{0j} + \sum_{i=1}^n x_i w_{ij}$$

Calculate output of the hidden unit by applying its activation functions over z_{inj}

$$z_j = f(z_{inj})$$

and send the output signal from the hidden unit to the input of output layer units.

Step 5: For each output unit $y_k(k = 1$ to $m)$, calculate the net input:

$$y_{ink} = w_{0k} + \sum_{j=1}^p z_j w_{jk}$$

and apply the activation function to compute output signal

$$y_k = f(y_{ink})$$

Back propagation of error (Phase II)

Step 6: Each output unit $y_k(k = 1$ to $m)$ receives a target pattern corresponding to the input training pattern and computes the error correction term:

$$\delta_k = (t_k - y_k) f'(y_{ink})$$

On the basis of the calculated error correction term, update the change in weights and bias:

$$\Delta w_{jk} = \alpha \delta_k z_j, \quad \Delta w_{0k} = \alpha \delta_k$$

Step 7: Each hidden unit ($z_j, j=1$ to p) sums its delta inputs from the output units:

$$\delta_{inj} = \sum_{k=1}^n \delta_k w_{jk}$$

On the basis of the calculated δ_j , update the change in weights and bias:

$$\Delta v_{ij} = \alpha \delta_j x_i, \quad \Delta v_{0j} = \alpha \delta_j$$

Weight and bias updation (Phase III):

Step 8: Each output unit ($y_k, k=1$ to m) updates the bias and weights:

$$w_{jk}(\text{new}) = w_{jk}(\text{old}) + \Delta w_{jk}$$

$$w_{0k}(\text{new}) = w_{0k}(\text{old}) + \Delta w_{0k}$$

Each hidden unit ($z_j, j=1$ to p) updates its bias and weights:

$$v_{ij}(\text{new}) = v_{ij}(\text{old}) + \Delta v_{ij}$$

$$v_{0j}(\text{new}) = v_{0j}(\text{old}) + \Delta v_{0j}$$

Step 9: Check for the stopping condition. The stopping condition may be certain number of epochs reached or when the actual output equals the target output.

The above algorithm uses the incremental approach for updation of weights, i.e., the weights are being changed immediately after a training pattern is presented. There is another way of training called **batch-mode training**, where the weights are changed only after all the training patterns are presented. The effectiveness of two approaches depends on the problem, but batch-mode training requires additional local storage

for each connection to maintain the immediate weight changes. When a BPN is used as a classifier, it is equivalent to the optimal Bayesian discriminant function for asymptotically large sets of statistically independent training patterns.