

Associative Learning

An association is any link between a system's input and output such that when a pattern A is presented to the system it will respond with pattern B.

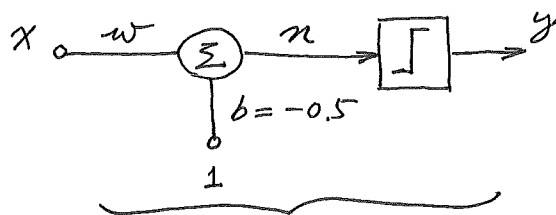
When two patterns are linked by an association, the input pattern is often referred to as stimulus. Likewise, the output pattern is referred to as the response.

Associations are so fundamental that they formed the foundation of the behaviorist school of psychology. This branch of psychology attempted to explain much of animal and human behavior by using associations and rules for learning associations.

One of the earliest influences on the behaviorist school of psychology was the classic experiment of Ivan Pavlov, in which he trained a dog to salivate at the sound of a bell, by ringing the bell whenever food was presented. This is an example of what is now called classical conditioning.

It was to provide a biological explanation for some of this behavior that led Donald Hebb to his postulate in *The Organization of Behavior*: The Hebbian learning.

Simple Associative Network



hard limit neuron

$$x = \begin{cases} 1, & \text{stimulus} \\ 0, & \text{no stimulus} \end{cases}$$

$$y = \begin{cases} 1, & \text{response} \\ 0, & \text{no response} \end{cases}$$

The presence of an association between the stimulus $x=1$, and the response $y=1$ is dictated by the value of w .

$$y = \text{hardlimit}(-wx + b) = \text{hardlimit}(-wx - 0.5)$$

The network will respond to the stimulus only if w is greater than $-b = 0.5$.

Two types of inputs:

Unconditioned stimulus: food

Unconditioned response: salivation

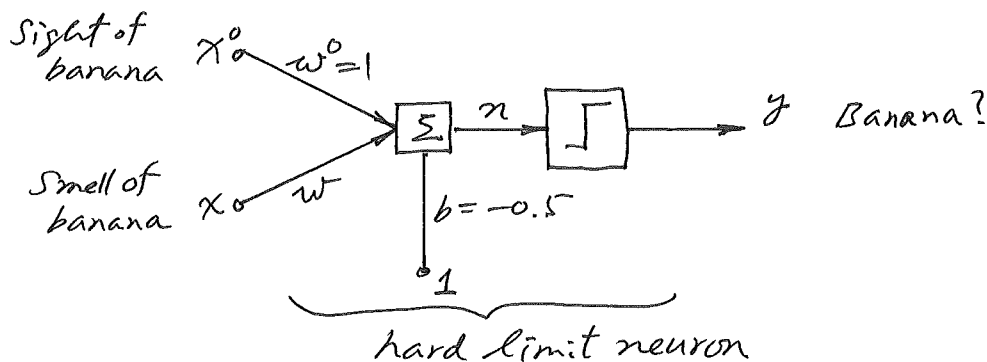
Conditioned stimulus: bell

Conditioned response: salivation

Banana Associator:

Unconditioned stimulus (banana shape): x^0

Conditioned stimulus (banana smell): x



$$x^0 = \begin{cases} 1, & \text{shape detected} \\ 0, & \text{" not detected} \end{cases}$$

$$x = \begin{cases} 1, & \text{smell detected} \\ 0, & \text{" not detected} \end{cases}$$

First, the network associates the shape of banana, but not the smell, with a response indicating the fruit is a banana.

$$y = \text{hardlimit}(w^0 x^0 + w x + b)$$

$$w = 0, \quad w^0(1) + b > 0 \Rightarrow w^0 > -b = 0.5$$

$$\text{Assign: } w^0 = 1, \quad w = 0$$

$$\Rightarrow y = \text{hardlim}(x^0 - 0.5)$$

Unsupervised Hebb Rule:

When should an association be learned? It is generally accepted that both animals and humans tend to associate things that occur simultaneously.

Paraphrasing Hebb:

If a banana smell stimulus occurs simultaneously with a banana concept response (activated by some other stimulus such as the sight of a banana shape), the network should strengthen the connection between them so that later it can activate its banana concept in response to the banana smell alone.

Following the Hebb rule:

$$w_{ji}(p) = w_{ji}(p-1) + \alpha y_j(p) x_i(p)$$

w_{ji} : weight between a neuron's input x_i and output y_j .

α : learning rate: dictates how many times a stimulus and response must occur together before an association is made.

In the banana associator, if we make

$$w(1) + b > 0 \Rightarrow w > -b = 0.5$$

an association is made, regardless of the value of x^0 .

Note: The above rule uses only signals available within the layer containing the weights being updated. Rules that satisfy this condition are called local learning rules.

This is in contrast to the backpropagation rule, in which the sensitivity (delta) must be propagated back from the final layer.

In vector form:

$$W(p) = W(p-1) + \alpha y(p) x^T(p)$$

Learning is performed in response to a series of inputs presented in time (the training sequence):

$$x(1), x(2), \dots, x(P)$$

Returning to the banana associator, with initial weights

$$w^0 = 1, \quad w(0) = 0.$$

Smell sensor works reliably, but the shape sensor operates only intermittently (on even time steps).

Thus, the training sequence will consist of the repetitions of the following two sets of inputs:

$$\{x^0(1) = 0, x(1) = 1\}, \{x^0(2) = 1, x(2) = 1\}, \dots$$

w^0 will remain constant

w will be updated at each iteration, using the unsupervised Hebb rule with a learning rate of 1:

$$w(p) = w(p-1) + y(p)x(p).$$

The output of the first iteration ($p=1$) is

$$\begin{aligned} y(1) &= \text{hardlim}(w^0 x^0 + w(0)x(1) - 0.5) \\ &= \text{hardlim}(1)(0) + (0)(1) - 0.5 \\ &= 0 \quad (\text{no response}) \end{aligned}$$

The smell alone did not generate a response. Without a response, the Hebb rule does not alter w :

$$w(1) = w(0) + y(1)x(1) = 0 + (0)(1) = 0$$

In the second iteration ($p=2$), both the banana's shape and smell are detected and the network responds:

$$\begin{aligned} y(2) &= \text{hardlim}(w^0 x^0(2) + w(1)x(2) - 0.5) \\ &= \text{hardlim}(1)(1) + (0)(1) - 0.5 \\ &= 1 \quad (\text{banana}) \end{aligned}$$

The weight changed is

$$w(2) = w(1) + y(2)x(2) = 0 + (1)(1) = 1$$

When the sight detector fails again, in the third iteration, the network responds anyway.

$$\begin{aligned} y(3) &= \text{hardlim}(-w^0 x^0(3) + w(2) x(3) - 0.5) \\ &= \text{hardlim}((1)(0) + (1)(1) - 0.5) \\ &= 1 \quad (\text{banana!}) \end{aligned}$$

$$w(3) = w(2) + y(3) x(3) = 1 + (1)(1) = 2$$

From now on, the network is capable of responding to bananas that are detected by either by sight or smell.

NOTE: 1. If we continue to present inputs and update w , the weight w will become arbitrarily large.

This is at odds with the biological systems that inspired the Hebb rule. Synapses cannot grow without bound.

2. There is no mechanism for weights to decrease. If the inputs or outputs experience any noise, every weight will grow (however slowly) until the network responds to any stimulus.

Hebb Rule with Decay:

One way to improve the Hebb rule is by adding a weight decay term:

$$\begin{aligned} w(p) &= w(p-1) + \alpha y(p) x^T(p) - \gamma w(p-1) \\ &= (1 - \gamma) w(p-1) + \alpha y(p) x^T(p), \end{aligned}$$

where, γ , the decay rate, is a positive constant less than one. As γ approaches zero, the learning law becomes the standard rule. As γ approaches one, the learning law quickly forgets old inputs and remembers only the most recent patterns.

This keeps the weight matrix from growing without bound.

The maximum weight value $w_{ji}^{\max}$ is determined by γ . To maximize learning, set both x_j and y_i to a value of 1 for all p and solve for the steady state weight (i.e., when both new and old weights are equal):

$$w_{ji} = (1 - \gamma) w_{ji} + \alpha y_j x_i$$

$$w_{ji} = (1 - \gamma) w_{ji} + \alpha$$

$$w_{ji} = \frac{\alpha}{\gamma}$$

Return to the banana associator. Use a decay rate γ of 0.1. The 1st iteration, with only the smell stimulus present, is the same:

$$y(1) = 0 \quad (\text{no response}), \quad w(1) = 0.$$

The 2nd iteration also produce the identical results:

$$y(2) = 1 \quad (\text{banana}), \quad w(2) = 1.$$

The 3rd iteration. The network has learned to respond to smell, and the weight continues to increase. However, this time the weight increases by only 0.9, instead of 1.0:

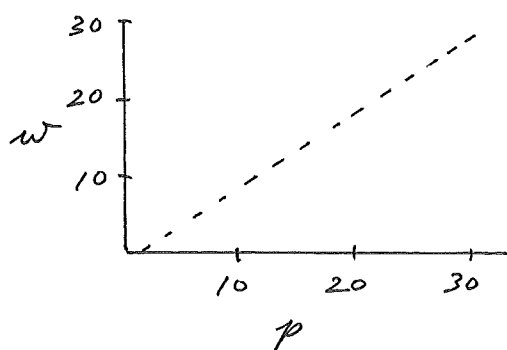
$$\begin{aligned} w(3) &= w(2) + y(3) x(3) - 0.1 w(2) \\ &= 1 + (1)(1) - 0.1(1) = 1.9 \end{aligned}$$

The decay term limits the weight's value, so that no matter how often the association is reinforced, w will never increase beyond $w^{\max}$

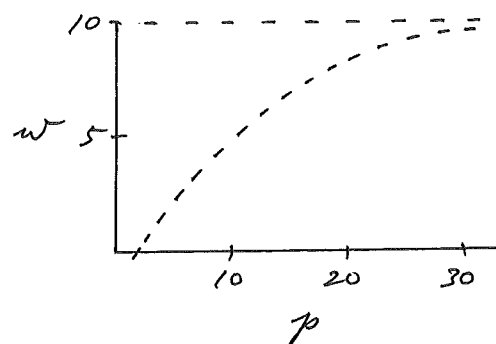
$$w^{\max} = \frac{\alpha}{\gamma} = \frac{1}{0.1} = 10$$

The new rule also ensures that associations learned by the network will not be artifacts of noise.

Any small random increases will soon decay away.



Response without decay



Response with decay

NOTE: The environment must be counted on to occasionally present all stimuli that have associations.

Without reinforcement, association will decay away:

For example, if $y_j = 0$,

$$w_{ji}(p) = (1 - \gamma) w_{ji}(p-1)$$

If $\gamma = 0.1$, this reduces to

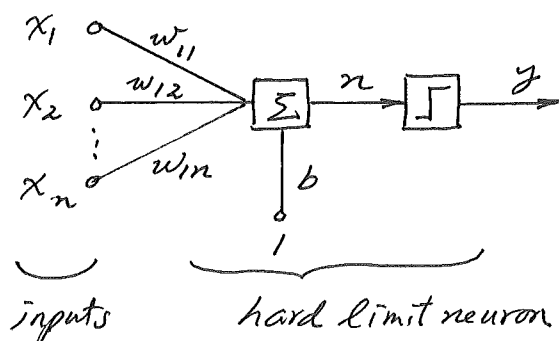
$$w_{ji}(p) = 0.9 w_{ji}(p-1)$$

Therefore, w_{ji} will be decreased by 10% at each presentation for which $y_j = 0$.

Any association that was previously learned will eventually be lost!

Simple Recognition Network

Instar : Neuron with a vector input.



The input/output expression for the instar is

$$y = \text{hardlim}(Wx + b) = \text{hardlim}(w^T x + b)$$

The instar will be active whenever the inner product between the weight vector and the input vector is greater than or equal to $-b$:

$$w^T x \geq -b.$$

or

$$w^T x = \|w\| \|x\| \cos \theta \geq -b.$$

where θ is the angle between two vectors.

If x and w have the same length, then the inner product will be largest when $x = w$. Therefore, the instar will be active when x is "close" to w .

By setting the bias b appropriately, we can select how close the input vector must be to the weight vector in order to activate the instar.

If we set

$$b = -\|w\| \|x\|,$$

then the instar will only be active when x points in exactly the same direction as w ($\theta = 0$). Thus, we will have a neuron that recognizes only the pattern w .

If we would like the instar to respond to any pattern near w (θ small), then we can increase b to some value larger than $-\|w\| \|x\|$.

Instar Rule:

One problem of the Hebb rule with decay was that it required stimuli to be repeated or associations would be lost.

A better rule might allow weight decay only when the instar is active ($y \neq 0$).

Consider the Hebb rule:

$$w_{ji}(p) = w_{ji}(p-1) + \alpha y_j(p) x_i(p)$$

Adding a decay term,

$$w_{ji}(p) = w_{ji}(p-1) + \alpha y_j(p) x_i(p) - \gamma y_j(p) w_{ji}^{\text{old}}$$

We can simplify by setting γ equal to α (so new weight values are learned at the same rate old values decay),

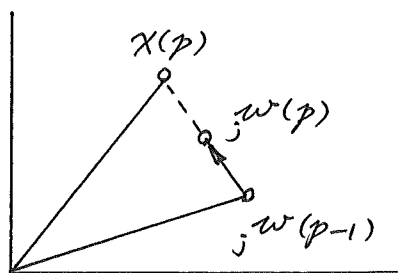
$$w_{ji}(p) = w_{ji}(p-1) + \alpha y_j(p) (x_i(p) - w_{ji}^{\text{old}})$$

In vector form,

$${}_jw(p) = {}_jw(p-1) + \alpha y_j(p) (x(p) - {}_jw(p-1))$$

Consider the case when the instar is active ($y_j = 1$).

$$\begin{aligned} {}_jw(p) &= {}_jw(p-1) + \alpha (x(p) - {}_jw(p-1)) \\ &= (1 - \alpha) {}_jw(p-1) + \alpha x(p). \end{aligned}$$

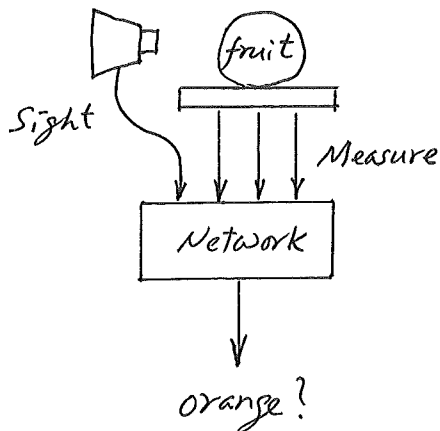


NOTE: When the instar is active, the weight vector is moved toward the input vector along a line between the old weight vector and the input vector.

If the input vectors are normalized, then ${}_jw$ will also be normalized once it has learned a particular vector x .

The distance the weight vector moves depends on the value of α .

Example:



Inputs:

visual input (orange): unconditioned stimulus

Measured inputs: conditioned stimulus.

$$x^0 = \begin{cases} 1, & \text{orange detected visually} \\ 0, & \text{orange not detected.} \end{cases}$$

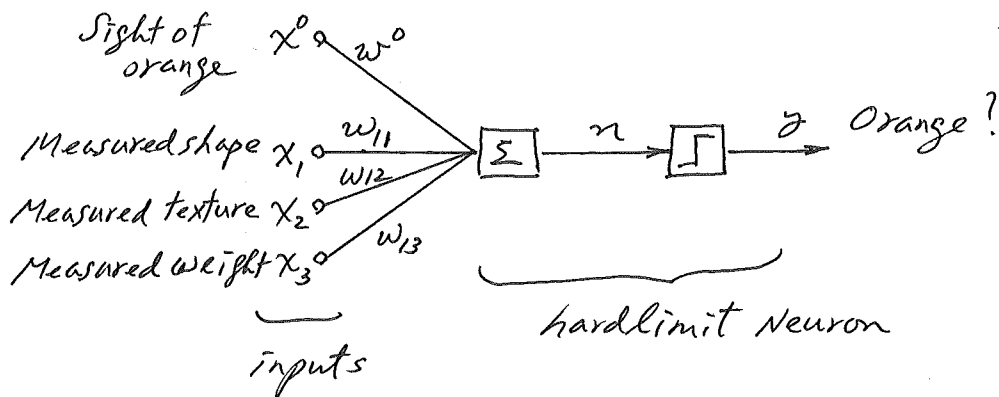
$$x = \begin{bmatrix} \text{shape} \\ \text{texture} \\ \text{weight} \end{bmatrix} = \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix}, \quad \|x\| = \sqrt{3}$$

The output of the network is

$$y = \text{hardlim}(w^0 x^0 + w x + b)$$

The bias: $b = -\|w\| \|x\| = -\|x\|^2 = -3$

Set bias as $b = -2$.



We need the network to have a constant association between the sight of an orange and its response, so w^0 will be set greater value than $-b$. But initially, the network should not respond to any combination of fruit measurements:

$$w^0 = 3 \quad w(0) = {}_1w^T(0) = [0 \ 0 \ 0]$$

Use the instar rule to update the weights with $d = 1$.

$${}_1w(p) = {}_1w(p-1) + z(p)(x(p) - {}_1w(p-1))$$

Training sequence: repeated presentation of an orange.
 Assume - visual system ^{only} operates correctly on even time steps.
 (due to a fault in its construction).

$$\left\{ x^0(1)=0, x(1)=\begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} \right\}, \left\{ x^0(2)=1, x(2)=\begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} \right\}, \dots$$

1st:

$$\begin{aligned} y(1) &= \text{hardlim}(w^0 p(1) + W x(1) - 2) \\ &= \text{hardlim}\left((3)(0) + [0 \ 0 \ 0] \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} - 2\right) = 0 \quad (\text{no response}) \end{aligned}$$

$${}_1w(1) = {}_1w(0) + y(1)(x(1) - {}_1w(0)) = 0 \quad (\text{not adjusted})$$

2nd:

$$\begin{aligned} y(2) &= \text{hardlim}(w^0 x(2) + W x(2) - 2) \\ &= \text{hardlim}\left((3)(1) + [0 \ 0 \ 0] \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} - 2\right) = 1 \quad (\text{orange}) \end{aligned}$$

$$\begin{aligned} {}_1w(2) &= {}_1w(1) + y(2)(x(2) - {}_1w(1)) \\ &= \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} + 1 \left(\begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} - \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \right) = \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} \end{aligned}$$

3rd:

$$\begin{aligned} y(3) &= \text{hardlim}(w^0 x^0(3) - W x(3) - 2) \\ &= \text{hardlim}\left((3)(0) + [1 \ -1 \ -1] \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} - 2\right) = 1 \quad (\text{orange}) \end{aligned}$$

$$\begin{aligned} {}_1w(3) &= {}_1w(2) + y(3)(x(3) - {}_1w(2)) \\ &= \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} + 1 \left(\begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} - \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} \right) = \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} \end{aligned}$$

Having completely learned the measurements, the weights stop changing. The network has learned to recognise an orange by its measurements, even when its visual detection system fails.

Kohonen Rule:

$$w_i(p) = w_i(p-1) + \alpha (x(p) - w_i(p-1)), \text{ for } i \in I(p).$$

Unlike the instar rule, learning is not proportional to the neuron's output $y_i(p)$. Instead, learning occurs when the neuron's index i is a member of the set $I(p)$.

Kohonen rule can be made equivalent to the instar rule by defining $I(p)$ as the set of all i such that

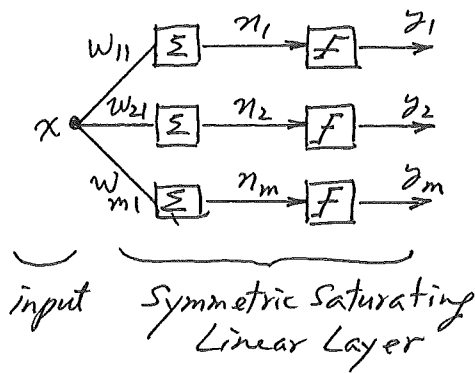
$$y_i(p) = 1.$$

The advantage of the Kohonen rule is that it can also be used with other definitions. It is useful for training networks such as the self-organizing feature map.

Simple Recall Network:

Outstar Network:

Has scalar input and a vector output.



The input/output association is

$$y = \text{sat}(\text{lin}(Wx))$$

: recall a vector containing values of -1 or 1.

If $x = 1$, stimulus, then

$$y = \text{sat}(\text{lin}(Wx)) = \text{sat}(\text{lin}(W \cdot 1)) = w$$

: w is an output vector associated with input $x = 1$.

Outstar: Column of the weight matrix is set to the desired vector.

Instar: Row of the weight matrix is set to the desired vector.

Outstar Rule:

Weight decay term of the Hebb rule:

proportional to the output of the network, y_j : instar

" to the input of the network, x_i : outstar

$$w_{ji}(p) = w_{ji}(p-1) + \alpha y_j(p) x_i(p) - \gamma x_i(p) w_{ji}(p-1)$$

Setting $\gamma = \alpha$,

$$w_{ji}(p) = w_{ji}(p-1) + \alpha (y_j(p) - w_{ji}(p-1)) x_i(p)$$

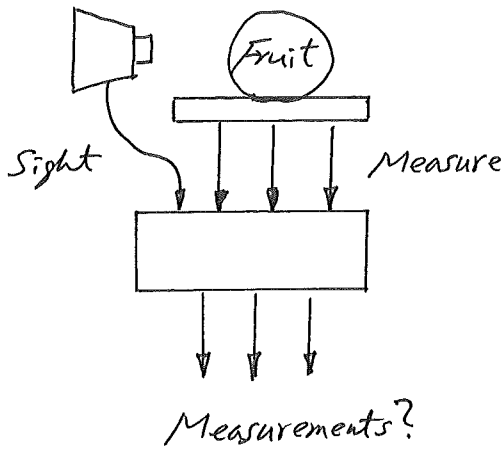
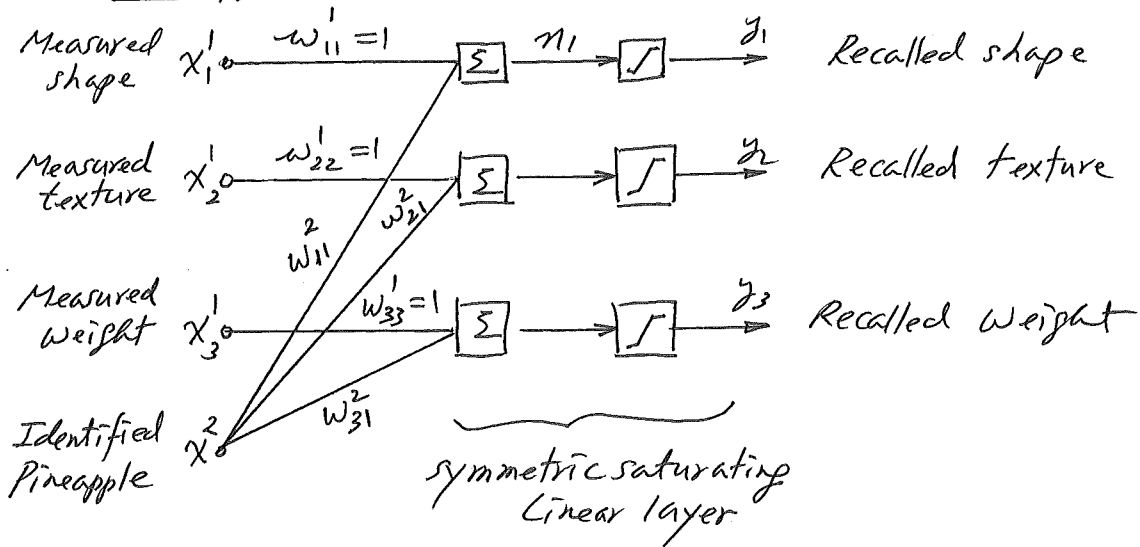
The outstar rule has properties complementary to the instar rule. Learning occurs whenever x_i is nonzero (instead of y_j). When learning occurs, column w_j moves toward the output vector.

In vector form,

$$w_j(p) = w_j(p-1) + \alpha (y(p) - w_j(p-1)) x_i(p)$$

where w_j is the j th column of the matrix W .

Pineapple Recaller:



The outputs of the network:

$$y = \text{satlin}(W^0 x^0 + W x)$$

$$W^0 = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}, \quad W(0) = 0, \quad \text{initial}$$

x^0 : unconditioned stimulus

x : conditioned stimulus

$$x = \begin{cases} 1, & \text{if a pineapple can be seen} \\ 0, & \text{otherwise} \end{cases}$$

$$x^0 = \begin{bmatrix} \text{shape} \\ \text{texture} \\ \text{weight} \end{bmatrix}$$

$$x^{\text{pineapple}} = \begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix}, \quad \text{measured values will only be available on even iterations.}$$

Training sequence:

$$\left\{ x^0(0) = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}, x(1) = 1 \right\}, \left\{ x^0(2) = \begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix}, x(2) = 1 \right\}, \dots$$

1st: $y(1) = \text{satlin}(W^0 x^0(1) + W x(1))$
 $= \text{satlin}\left(\begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} + \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} 1\right) = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}, \text{ no response}$

$$w_j(p) = w_j(p-1) + (y(p) - w_j(p-1)) x(p)$$

$$w_1(1) = w_1(0) + (y(1) - w_1(0)) x(1)$$

$$= \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} + \left(\begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} - \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \right) 1 = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \quad \text{no changes}$$

2nd: $y(2) = \text{satlin}\left(\begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix} + \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} 1\right) = \begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix}$ measurements given

$$w_1(2) = w_1(1) + (y(2) - w_1(1)) x(2)$$

$$= \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} + \left(\begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix} + \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \right) 1 = \begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix}$$

3rd: Measurements are unavailable,

$$y(3) = \text{satlin}\left(\begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} + \begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix} 1\right) = \begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix} \quad \text{measurements recalled.}$$

$$w_1(3) = w_1(2) + (y(3) - w_1(2)) y(2)$$

$$= \begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix} + \left(\begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix} - \begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix} \right) 1 = \begin{bmatrix} -1 \\ -1 \\ 1 \end{bmatrix}$$

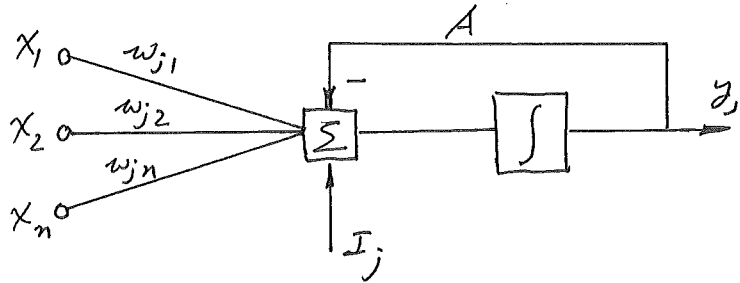
From now on, the weights will no longer change values unless a pineapple is seen with different measurements.

Grossberg Learning

Mathematically formulate Hebb's law in continuous-time.

$$\frac{\partial y_j(t)}{\partial t} = -A y_j(t) + I_j(t) + \sum_{i=1}^n w_{ji} y_i(t)$$

$\parallel$
 x_i



$y_j(t)$: output of the neuron

$I_j(t)$: external input to the neuron

$w_{ji}(t)$: weight connecting the output of i th neuron to the input of the j th neuron

A : positive constant controlling the decay of the output in the absence of any other inputs.

In Pavlov's dog:

I_j : plate of food

x_i : bell ringing

y_j : dog is salivating

A : dog will eventually discontinue salivating if it is not sufficiently stimulated by either the plate of food or the bell.

Learning of the network:

$$\frac{\partial w_{ji}(t)}{\partial t} = -F w_{ji}(t) + G y_j(t) y_i(t)$$

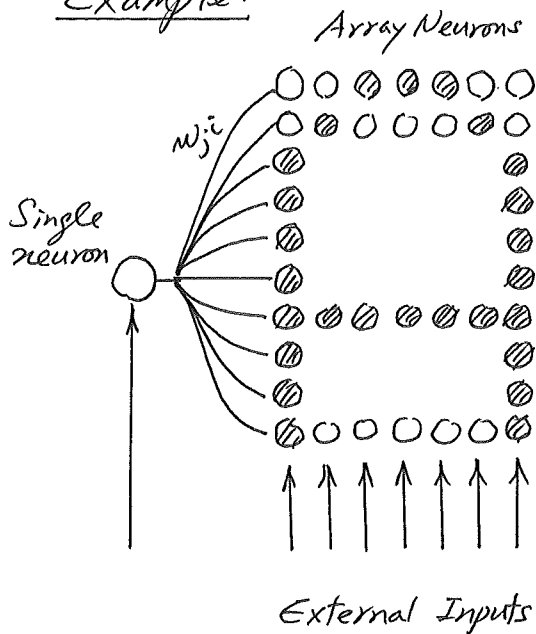
where F : forgetting rate (much smaller than A)

G : learning rate in Hebb rule.

Time constants for A : in the order of milliseconds
 " F : " seconds

e.g., $e^{-At} = e^{-t/T}$, $T = 10^{-2} \Rightarrow A = 100$
 $e^{-Ft} = e^{-t/T}$, $T = 1 \Rightarrow F = 1$

Example:



Training:

External inputs are used to stimulate the single neuron and to stimulate the array with a pattern (A, for example)

The array neurons will adjust their output $y_j(t)$ and weight $w_{ji}(t)$ from the single neuron.

Array neurons which are not stimulated by an external input (O in the figure) have a small output and their weight slowly decay.

$\Rightarrow$ Array neurons do not respond to the single neuron's stimulation.

On the other hand, array neurons stimulated by an external input have large outputs and increase the strength of the weight, learning to respond to the stimulus of the single neuron.

If the above external stimulations are repeated a sufficient number of times, the array neurons eventually learn to produce the pattern when only the single neuron stimulates them.

Network is trained to respond to the conditioned stimulus in the absence of the unconditioned stimulus.

By adding additional single neurons, and retraining the array neurons to respond with a different pattern (say, B), the network can be made useful.