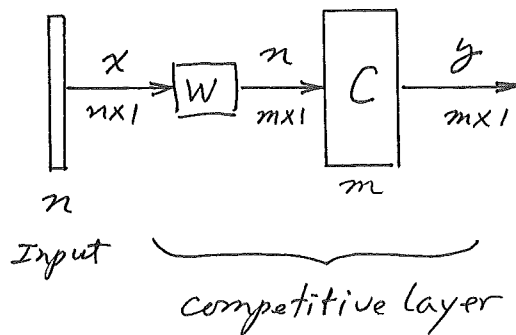


Competitive Networks

Competitive Layer:



$$y = \text{compet}(n)$$

: find the index i^* of the neuron with the largest net input, and setting its output to 1 (with ties going to the neuron with the lowest index).

All other outputs are set to 0.

$$y_i = \begin{cases} 1, & i = i^* \\ 0, & i \neq i^* \end{cases}, \text{ where } n_{i^*} \geq n_i, \forall i, \text{ and } i^* \leq \bar{i}, \forall n_i = n_{i^*}$$

The prototype vectors are stored in the rows of W . The net input n calculates the distance between the input vector x and each prototype $i^T w$ (assuming vectors have normalized lengths of L).

The net input n_i of each neuron i is proportional to the angle θ_i between x and the prototype $i^T w$:

$$n = Wx = \begin{bmatrix} 1^T w^T \\ 2^T w^T \\ \vdots \\ m^T w^T \end{bmatrix} x = \begin{bmatrix} 1^T w^T x \\ 2^T w^T x \\ \vdots \\ m^T w^T x \end{bmatrix} = \begin{bmatrix} L^2 \cos \theta_1 \\ L^2 \cos \theta_2 \\ \vdots \\ L^2 \cos \theta_m \end{bmatrix}$$

The competitive layer assigns an output of 1 to the neuron whose weight vector points in the direction closest to the input vector.

Competitive Learning:

We can now design a competitive network classifier by setting the rows of W to the desired prototype vectors. However, we would like to have a learning rule that could be used to train the weights in a competitive network, without knowing the prototype vectors.

Recall the instar rule:

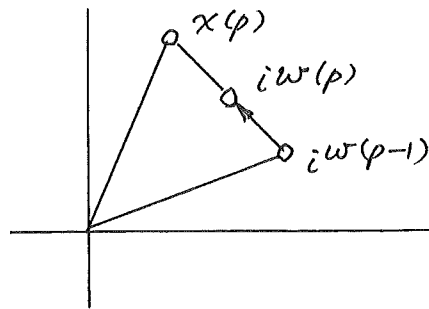
$$i w(p) = i w(p-1) + \alpha y_i(p) (x(p) - i w(p-1))$$

For the competitive network, y is only nonzero for the winning neuron ($i=i^*$). Therefore, we can get the same results using the Kohonen rule.

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

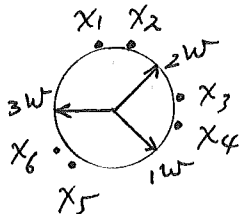
and

$$i w(p) = i w(p-1) \quad i \neq i^*$$



The row of the weight matrix that is closest to the input vector (or has the largest inner product with the input vector) moves toward the input vector.

Example:



Inputs:

$$\begin{aligned} x_1 &= \begin{bmatrix} -0.1861 \\ 0.9806 \end{bmatrix}, & x_2 &= \begin{bmatrix} 0.1861 \\ 0.9806 \end{bmatrix}, & x_3 &= \begin{bmatrix} 0.9806 \\ 0.1861 \end{bmatrix} \\ x_4 &= \begin{bmatrix} 0.9806 \\ -0.1861 \end{bmatrix}, & x_5 &= \begin{bmatrix} -0.5812 \\ -0.8137 \end{bmatrix}, & x_6 &= \begin{bmatrix} -0.8137 \\ -0.5812 \end{bmatrix} \end{aligned}$$

Initial weights, normalized, random

$$1w = \begin{bmatrix} 0.7071 \\ -0.7071 \end{bmatrix}, \quad 2w = \begin{bmatrix} 0.7071 \\ 0.7071 \end{bmatrix}, \quad 3w = \begin{bmatrix} -1.0000 \\ 0.0000 \end{bmatrix}, \quad W = \begin{bmatrix} 1w^T \\ 2w^T \\ 3w^T \end{bmatrix}$$

For input x_2 :

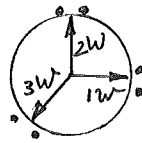
$$\begin{aligned} y &= \text{compet}(W x_2) = \text{compet} \left(\begin{bmatrix} 0.7071 & -0.7071 \\ 0.7071 & 0.7071 \\ -1.0000 & 0.0000 \end{bmatrix} \begin{bmatrix} 0.1861 \\ 0.9806 \end{bmatrix} \right) \\ &= \text{compet} \left(\begin{bmatrix} -0.5547 \\ 0.8321 \\ -0.1861 \end{bmatrix} \right) = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} \end{aligned}$$

The second neuron's weight vector ($2w$) was closest to x_2 , so it won the competition ($i^*=2$) and output is 1.

Now we apply the Kohonen learning rule to the winning neuron with $\alpha = 0.5$:

$$\begin{aligned} \vec{w}^{\text{new}} &= \vec{w}^{\text{old}} + \alpha (\vec{x}_2 - \vec{w}^{\text{old}}) \\ &= \begin{bmatrix} 0.7071 \\ 0.7071 \end{bmatrix} + 0.5 \left(\begin{bmatrix} 0.1861 \\ 0.8806 \end{bmatrix} - \begin{bmatrix} 0.7071 \\ 0.7071 \end{bmatrix} \right) = \begin{bmatrix} 0.4516 \\ 0.8438 \end{bmatrix} \end{aligned}$$

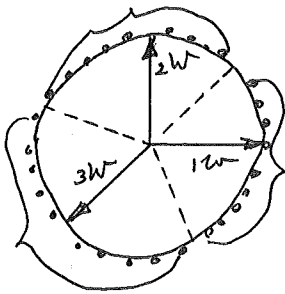
The Kohonen rule moves $\vec{w}$ closer to $\vec{x}_2$.



If we continue choosing input vectors at random and presenting them to the network, then at each iteration the weight vector closest to the input vector will move toward that vector.

Eventually, each weight vector will point at a different cluster of input vectors.

Each weight vector becomes a prototype for a different cluster.



Once the network has learned to cluster the input vectors, it will classify new vectors accordingly.

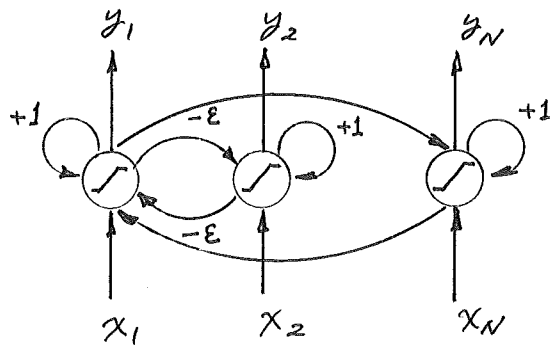
Problems with Competitive Layers:

Choice of learning rate: trade-off between the speed of learning and the stability of the final weight vectors.

$\alpha \sim 0$: slow learning, once a weight vector reaches the center of cluster, it will tend to stay close to the center.

$\alpha \sim 1$: fast learning, but it will continue to oscillate as different vectors in the cluster are presented.

Winner-Take-All Networks



The typical competitive neural network consists of a layer of processing elements (PEs), all receiving the same input. The PE with the best output (either maximum or minimum, depending on the criteria) will be declared the winner.

In digital computers, choosing a winner is incredibly simple (just a search for the largest value).

The concept of choosing a winner, however, often requires a global controller that compares each output with all the others, which is troublesome in distributed systems.

For this and other reasons (e.g., biological plausibility) we wish to construct a network that will find the largest (or smallest) output without global control.

We call this simple system a winner-take-all network.

Inputs: $x_1, x_2, \dots, x_N$

Outputs: $y_1, y_2, \dots, y_N$

$$y_k = \begin{cases} 1 & x_k \text{ largest} \\ 0 & \text{otherwise} \end{cases}$$

PE: has a semilinear nonlinearity (clipped linear region), has a self-excitatory connection with a fixed weight of +1,

(laterally connected to all other PEs by negative weights $-\epsilon$ (lateral inhibition), $0 < \epsilon < \frac{1}{N}$).

$$x_i \geq 0$$

Initial condition: zero output for no input.

Input is presented as an initial condition, i.e., it is presented for one sample and then pulled back.

The lateral inhibition drives all the outputs toward zero at an exponential rate.

As all the smaller PEs approach zero, the largest PE (PE k) will be less and less affected by the lateral inhibition.

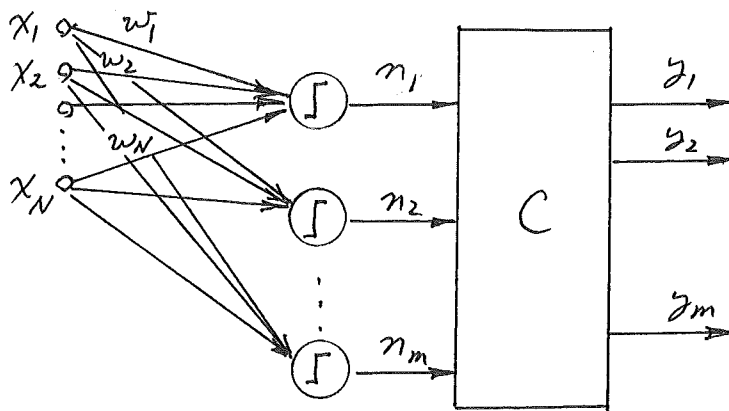
At this time the self-excitation drives its output high, which enforces the zero values of the other PEs.

Hence this solution is stable. Note that the PEs compete for the output activity, and only one wins it.

This network has feedback among the PEs. The output takes some time to stabilize, unlike the feedforward network, which are instantaneous.

Application: The amplitude difference at the input could be small, but at the output it is very clear, so the network amplified the differences, creating a selector mechanism that can be applied to many different applications.

A typical application of the winner-take-all network is at the output of another network (such as a linear associative memory (LAM)) to select the highest output. Thus the net makes a "decision" based on the most probable answer.



Clustering

The competitive rule allows a single-layer linear network to group and represent data samples that lie in a neighborhood of the input space. Each neighborhood is represented by a single output PE. This operation is commonly called clustering in pattern recognition.

From the point of view of the input space, clustering is dividing the space into local regions, each of which is associated with an output PE. The input space is divided as honeycomb. The weights of each PE represents points in the input space called prototype vectors.

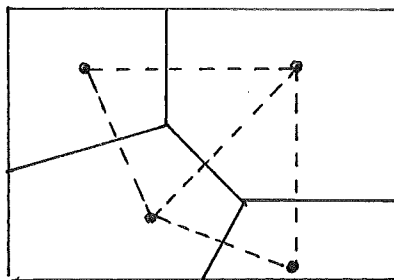
If we join prototype vectors by a line, its perpendicular bisector will meet other other bisectors, forming a division resembles a honeycomb. Mathematically this division is called Voronoi Tessellation, or simply a tessellation.

Data samples that fall inside the regions are assigned to the corresponding prototype vector. Clustering is therefore a continuous-to-discrete transformation.

The most important engineering application of competitive learning is vector quantization. In telecommunications, data reduction is needed for economic reasons.

With vector quantization, instead of transmitting the values of each data sample, the data is first categorized in clusters for which the centers, called codebook entries, are known to the transmitter and the receiver. Then just the cluster number is transmitted instead of the data samples.

At the receiver the cluster number is replaced with the codebook entry to recreate the transmitted signal.



The ultimate requirement of a vector quantizer is to have a set of clusters that minimizes the distance between the centers of each cluster and the input that falls into each cluster.

K-means clustering algorithm: (Duda & Hart 1973)

To find the best division of N samples by K clusters C_i such that the total distance between the clustered samples and their respective centers (i.e., the total variance) is minimized:

$$J = \sum_{i=1}^K \sum_{n \in C_i} |x_n - \gamma_i|^2$$

where γ_i is the center of class i .

- Steps:
1. randomly assign samples to the class C_i ,
 2. compute the centers according to

$$\gamma_i = \frac{1}{N_i} \sum_{n \in C_i} x_n$$

3. reassign the samples to the nearest cluster.
4. reiterate

Gradient estimate:

$$\Delta \gamma_i(n) = \gamma (x(n) - \gamma_i(n))$$

Recall the competitive rule:

$$w_{i^*}(n+1) = w_{i^*}(n) + \gamma (x(n) - w_{i^*}(n))$$

where i^* is the PE that wins the competition.
All other PEs keep their previous weights.

The gradient estimate is exactly the competitive update if we link the cluster center with the j th PE weight.

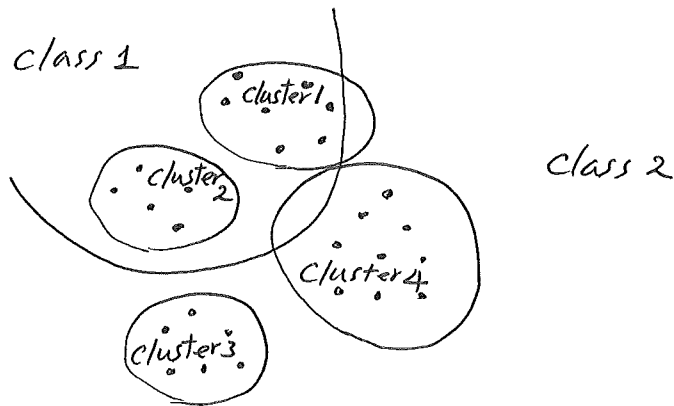
Competitive networks thus implement an on-line version of K-means clustering instead of the required batch adaptation of K-means.

Clustering and Classification

Clustering is the process of grouping input samples that are spatial neighbors.

Classification involves the labeling of input samples via some external criterion.

clustering is an unsupervised process of grouping,
Classification is supervised.



Data from each class tends to be dense, and there is a natural valley between classes.

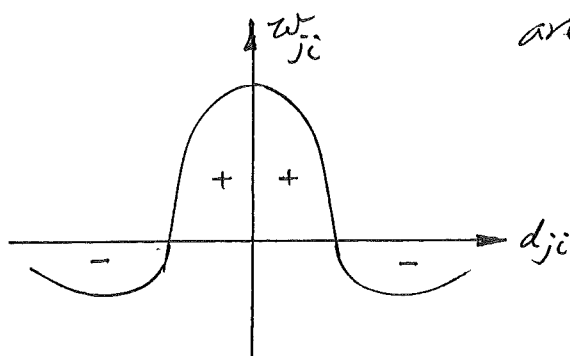
In such cases, clustering can be a preprocessor for classification.

Soft Competition:

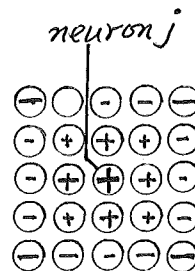
Hard competition: There is only one winner

Soft competition: Not only the winner but also its neighbors are active.

Creates a "bubble" of activity in the output space where the closest PE is the most active (highest output) and its neighbors are less active.

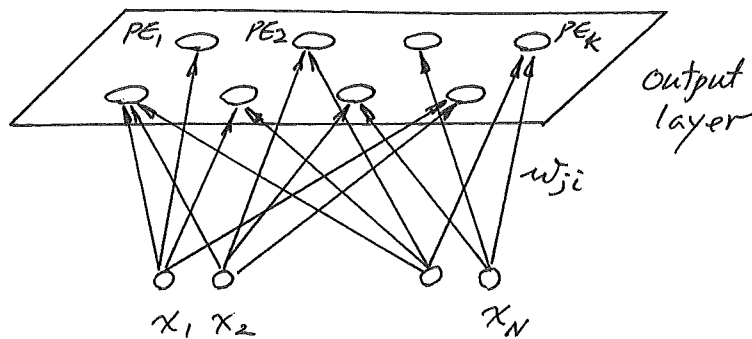


"Mexican Hat" distribution



ON-center/ OFF-surround

Self-Organizing (Feature) Map



The Kohonen self-organizing map (SOM) network performs a mapping from a continuous input space to a discrete output space, preserving the topological properties of the input. This means that points close to each other in the input space are mapped to the same or neighboring PEs in the output space. The basis of the Kohonen SOM network is soft competition among the PEs in the output space.

The Kohonen SOM is a fully connected, single-layer linear network. The output is generally organized in a one- or two-dimensional arrangement of PEs, which are called neighborhood.

The SOM network first determines the winning neuron i^* , then the weight vectors for all neurons within a certain neighborhood of the winning neuron are updated using the Kohonen rule,

$$w_i(n+1) = w_i(n) + \Lambda_{i,i^*} \eta(n) (x(n) - w_i(n))$$

where Λ_{i,i^*} is a neighborhood function centered at the winning neuron i^* . Typically, both the neighborhood and the step size change with the iteration number.

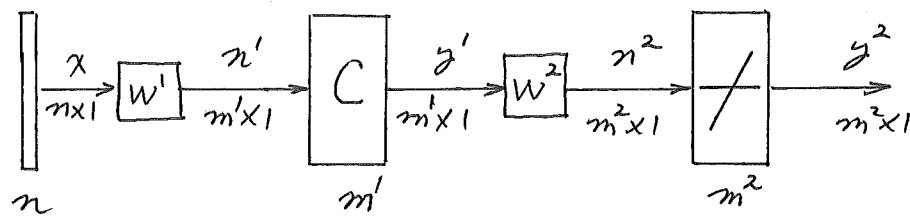
The neighborhood function Λ is normally a Gaussian:

$$\Lambda_{i,i^*}(n) = \exp\left(\frac{-d_{i,i^*}^2}{2\sigma^2(n)}\right)$$

with a variance that decreases with iteration.

Learning Vector Quantization (LVQ)

: Creating classifiers from competitive networks



$$n'_i = -\|w^1_i - x\|$$

$$y^2 = W^2 y'$$

$$y' = \text{compet}(n')$$

Competitive layer

Linear layer

Unsupervised

Supervised

In the LVQ network, each neuron in the first layer is assigned to a class, with several neurons often assigned to the same class. Each class is then assigned to one neuron in the second layer.

The winning neuron indicates a subclass, rather than a class. There may be several different neurons (subclasses) that make up each class.

The second layer of the LVQ network is used to combine subclasses into a single class.

$$W^2 = \begin{bmatrix} 0 \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix} \leftarrow \text{class : row } k$$

↑
subclass : column i

$$w^2_{ki} = 1 \quad : \quad \text{subclass } i \text{ is a part of class } k$$

LVQ Learning:

Combines the competitive learning with supervision.

Set of examples (training set):

$$\{x_1, t_1\}, \{x_2, t_2\}, \dots, \{x_n, t_n\}$$

Each target vector must contain only zeros, except for a single 1. The row in which the 1 appears indicates the class to which the input vector belongs.

Define W^2 : If hidden neuron i is to be assigned to class k , then set $w_{ki}^2 = 1$.

Once W^2 is defined, it will never be altered. The hidden weights W^1 are trained with a variation of the Kohonen rule:

At each iteration, an input vector x is presented to the network, and the distance from x to each prototype vector is computed. The hidden neurons compete, neuron i^* wins the competition, and the i^* th element of y^1 is set to 1.

Next, y^1 is multiplied by W^2 to get the final output y^2 , which also has one nonzero element, k^* , indicating that x is being assigned to class k^* .

The Kohonen rule is used to improve the hidden layer of the LVQ network in two ways.

First, if x is classified correctly, then we move the weights i^*w^1 of the winning hidden neuron toward x .

$$i^*w^1(n) = i^*w^1(n-1) + \alpha(x(n) - i^*w^1(n-1)),$$

$$\text{if } y_{k^*}^2 = t_{k^*} = 1.$$

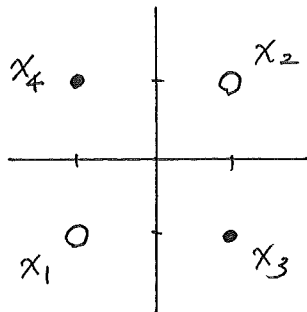
Second, if x was classified incorrectly, then we know that the wrong hidden neuron won the competition, and therefore we move its weights i^*w^1 away from x .

$$i^*w^1(n) = i^*w^1(n-1) - \alpha (x(n) - i^*w^1(n-1)),$$

$$\text{if } y_{k^*}^2 = 1 \neq t_{k^*} = 0.$$

The result will be that each hidden neuron moves toward vectors that fall into the class for which it forms a subclass and away from vectors that fall into other classes.

Example:



We want to train an LVQ network to solve the following classification problem:

$$\text{Class 1: } \{x_1 = \begin{bmatrix} -1 \\ -1 \end{bmatrix}, x_2 = \begin{bmatrix} 1 \\ 1 \end{bmatrix}\}$$

$$\text{Class 2: } \{x_3 = \begin{bmatrix} 1 \\ -1 \end{bmatrix}, x_4 = \begin{bmatrix} -1 \\ 1 \end{bmatrix}\}$$

Training sets:

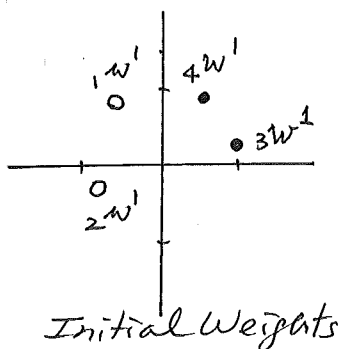
$$\{x_1 = \begin{bmatrix} -1 \\ -1 \end{bmatrix}, t_1 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}\}, \{x_2 = \begin{bmatrix} 1 \\ 1 \end{bmatrix}, t_2 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}\}$$

$$\{x_3 = \begin{bmatrix} 1 \\ -1 \end{bmatrix}, t_3 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}\}, \{x_4 = \begin{bmatrix} -1 \\ 1 \end{bmatrix}, t_4 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}\}$$

The output layer weight matrix:

$$W^2 = \begin{bmatrix} 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \end{bmatrix}$$

W^2 connects hidden neurons 1 and 2 to output neuron 1,
 " 3 and 4 " 2



The hidden layer weight matrix: random

$$1w^1 = \begin{bmatrix} -0.543 \\ 0.840 \end{bmatrix}, 2w^1 = \begin{bmatrix} -0.969 \\ -0.249 \end{bmatrix}, 3w^1 = \begin{bmatrix} 0.997 \\ 0.094 \end{bmatrix}, 4w^1 = \begin{bmatrix} 0.456 \\ 0.854 \end{bmatrix}$$

Initial Weights

Present input x_3 :

$$y' = \text{compet}(x') = \text{compet} \left(\begin{bmatrix} -\|w^1 - x_3\| \\ -\|2w^1 - x_3\| \\ -\|3w^1 - x_3\| \\ -\|4w^1 - x_3\| \end{bmatrix} \right)$$

$$= \text{compet} \left(\begin{bmatrix} -\|[-0.543 \ 0.840]^T - [1 \ -1]^T\| \\ -\|[-0.568 \ -0.248]^T - [1 \ -1]^T\| \\ -\|[0.997 \ 0.084]^T - [1 \ -1]^T\| \\ -\|[0.456 \ 0.554]^T - [1 \ -1]^T\| \end{bmatrix} \right) = \text{compet} \left(\begin{bmatrix} -2.40 \\ -2.11 \\ -1.08 \\ -2.03 \end{bmatrix} \right) = \begin{bmatrix} 0 \\ 0 \\ 1 \\ 0 \end{bmatrix}$$

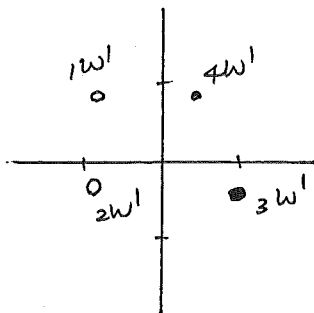
The 3rd hidden neuron has the closest weight vector to x_3 .
In order to determine which class this neuron belongs to, we multiply y' by w^2 ,

$$y^2 = w^2 y' = \begin{bmatrix} 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \end{bmatrix} \begin{bmatrix} 0 \\ 0 \\ 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$$

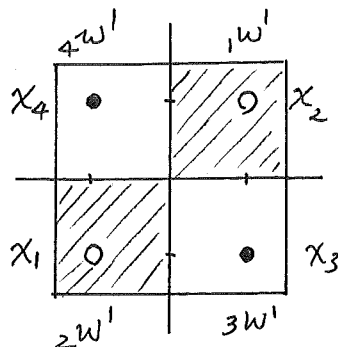
This output indicates that x_3 is a member of class 2. This is correct, so $3w^1$ is updated by moving it toward x_3 ,

$$3w^1(1) = 3w^1(0) + \alpha (x_3 - 3w^1(0))$$

$$= \begin{bmatrix} 0.997 \\ 0.084 \end{bmatrix} + 0.5 \left(\begin{bmatrix} 1 \\ -1 \end{bmatrix} - \begin{bmatrix} 0.997 \\ 0.084 \end{bmatrix} \right) = \begin{bmatrix} 0.998 \\ -0.453 \end{bmatrix}$$



After 1st iteration



After convergence