

Linear classifiers

Machine Learning

Hamid R Rabiee – Zahra Dehghanian
Spring 2025



Sharif University
of Technology

Classification problem

▶ Given: Training set

- ▶ labeled set of N input-output pairs $D = \{(\mathbf{x}^{(i)}, y^{(i)})\}_{i=1}^N$
- ▶ $y \in \{1, \dots, K\}$

▶ Goal: Given an input $\mathbf{x}$, assign it to one of K classes

▶ Examples:

- ▶ Spam filter
- ▶ Handwritten digit recognition

Linear classifiers

- ▶ Decision boundaries are linear in $\mathbf{x}$, or linear in some given set of functions of $\mathbf{x}$
- ▶ Linearly separable data: data points that can be exactly classified by a linear decision surface.
- ▶ Why linear classifier?
 - ▶ Even when they are not optimal, we can use their simplicity
 - ▶ are relatively easy to compute
 - ▶ In the absence of information suggesting otherwise, linear classifiers are an attractive candidates for initial, trial classifiers.

Two Category

➤ $f(\mathbf{x}; \mathbf{w}) = \mathbf{w}^T \mathbf{x} + w_0 = w_0 + w_1 x_1 + \dots w_d x_d$

- ▶ $\mathbf{x} = [x_1 \ x_2 \ \dots x_d]$
- ▶ $\mathbf{w} = [w_1 \ w_2 \ \dots w_d]$
- ▶ w_0 : bias

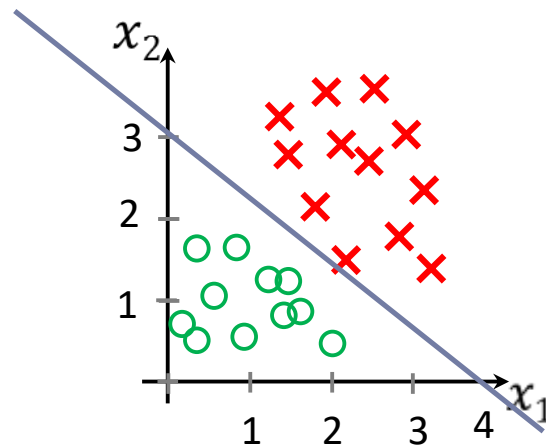
- ▶ if $\mathbf{w}^T \mathbf{x} + w_0 \geq 0$ then $\mathcal{C}_1$
- ▶ else $\mathcal{C}_2$

Decision surface (boundary): $\mathbf{w}^T \mathbf{x} + w_0 = 0$

$\mathbf{w}$ is orthogonal to every vector lying within the decision surface

Example

$$3 - \frac{3}{4}x_1 - x_2 = 0$$



if $\mathbf{w}^T \mathbf{x} + w_0 \geq 0$ then $\mathcal{C}_1$
else $\mathcal{C}_2$

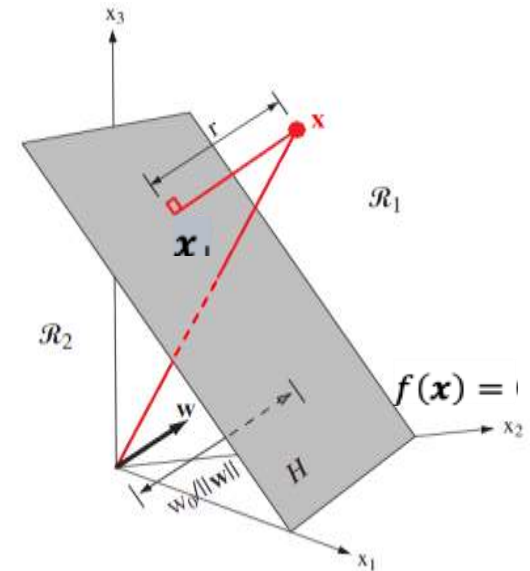
Linear classifier: Two Category

- Decision boundary is a $(d - 1)$ -dimensional hyperplane H in the d -dimensional feature space
 - The orientation of H is determined by the normal vector $[w_1, \dots, w_d]$
 - w_0 determines the location of the surface.
 - The normal distance from the origin to the decision surface is $\frac{w_0}{\|w\|}$

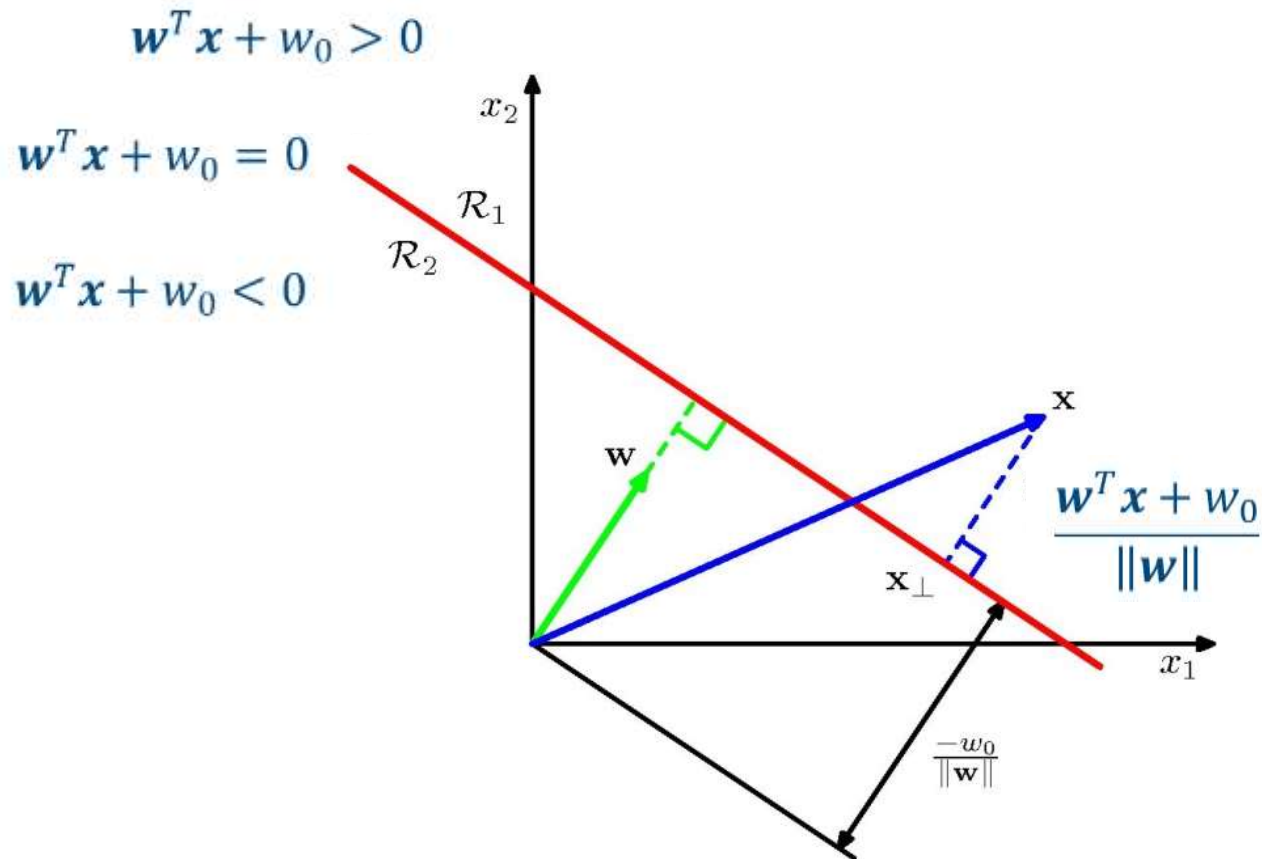
$$\mathbf{x} = \mathbf{x}_\perp + r \frac{\mathbf{w}}{\|\mathbf{w}\|}$$

$$\mathbf{w}^T \mathbf{x} + w_0 = r \|\mathbf{w}\| \Rightarrow r = \frac{\mathbf{w}^T \mathbf{x} + w_0}{\|\mathbf{w}\|}$$

gives a signed measure of the perpendicular distance r of the point $\mathbf{x}$ from the decision surface

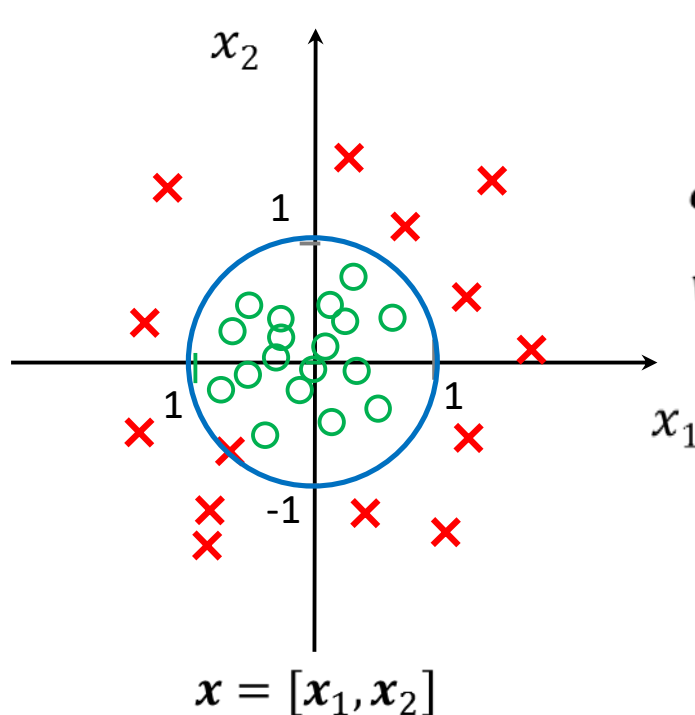


Linear boundary: geometry



Non-linear decision boundary

- ▶ Choose non-linear features
- ▶ Classifier still linear in parameters w



$$-1 + x_1^2 + x_2^2 = 0$$

$$\phi(x) = [1, x_1, x_2, x_1^2, x_2^2, x_1 x_2]$$

$$w = [w_0, w_1, \dots, w_m] = [-1, 0, 0, 1, 1, 0]$$

if $w^T \phi(x) \geq 0$ then $y = 1$
else $y = -1$

Cost Function for linear classification

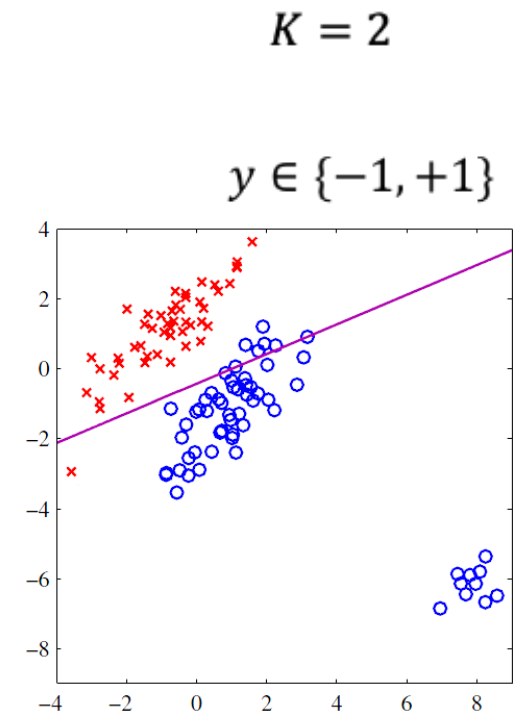
- Finding linear classifiers can be formulated as an optimization problem:
 - ▶ Select how to measure the prediction loss
 - ▶ Based on the training set $D = \{(\mathbf{x}^{(i)}, y^{(i)})\}_{i=1}^n$, a cost function $J(\mathbf{w})$ is defined
 - ▶ Solve the resulting optimization problem to find parameters:
 - ▶ Find optimal $\hat{f}(\mathbf{x}) = f(\mathbf{x}; \hat{\mathbf{w}})$ where $\hat{\mathbf{w}} = \underset{\mathbf{w}}{\operatorname{argmin}} J(\mathbf{w})$
- ▶ Criterion or cost functions for classification:
 - ▶ We will investigate several cost functions for the classification problem

SSE cost function for classification

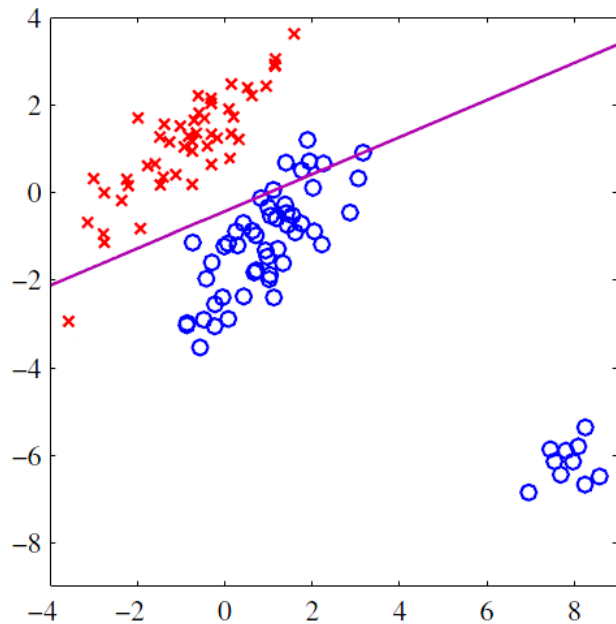
SSE cost function is not suitable for classification:

- ▶ Least square loss penalizes 'too correct' predictions (that they lie a long way on the correct side of the decision)
- ▶ Least square loss also lack robustness to noise

$$J(\mathbf{w}) = \sum_{i=1}^N (\mathbf{w}^T \mathbf{x}^{(i)} - y^{(i)})^2$$

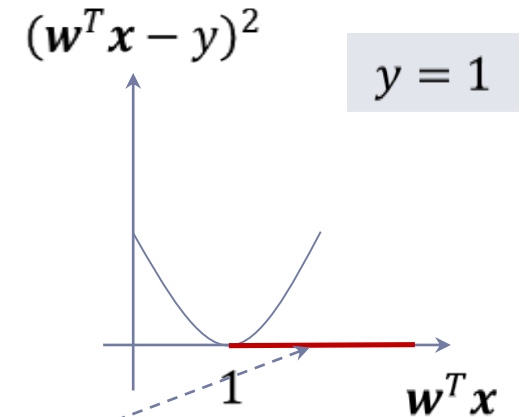


SSE cost function for classification

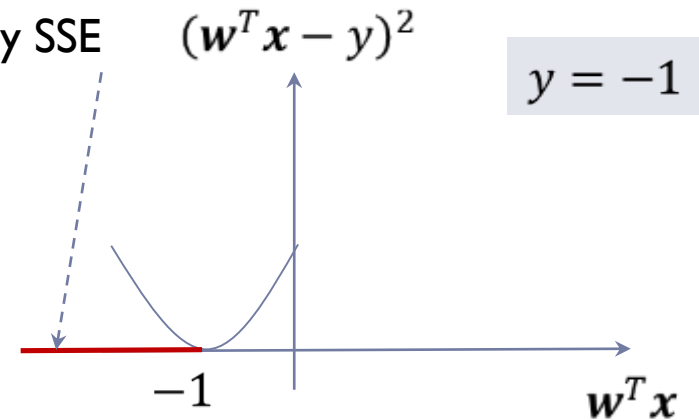


[Bishop]

$K = 2$



Correct predictions that
are penalized by SSE

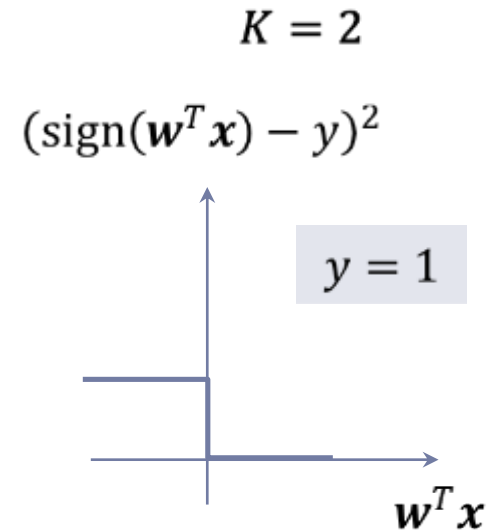


SSE cost function for classification

► Is it more suitable if we set $f(\mathbf{x}; \mathbf{w}) = g(\mathbf{w}^T \mathbf{x})$?

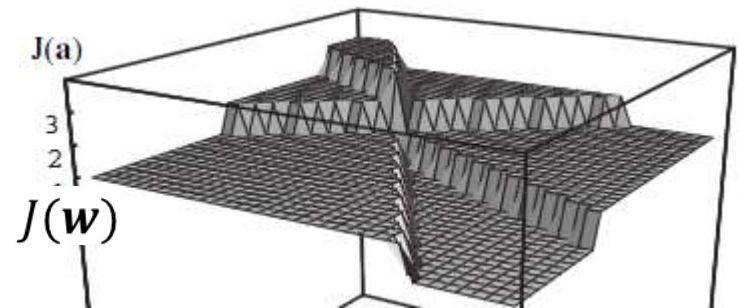
$$J(\mathbf{w}) = \sum_{i=1}^N (\text{sign}(\mathbf{w}^T \mathbf{x}^{(i)}) - y^{(i)})^2$$

$$\text{sign}(z) = \begin{cases} -1, & z < 0 \\ 1, & z \geq 0 \end{cases}$$



► $J(\mathbf{w})$ is a piecewise constant function shows the number of misclassifications

Training error incurred in classifying training samples



Perceptron algorithm

- ▶ Linear classifier
- ▶ Two-class: $y \in \{-1, 1\}$
 - ▶ $y = -1$ for C_2 , $y = 1$ for C_1
- ▶ Goal: $\forall i, \mathbf{x}^{(i)} \in C_1 \Rightarrow \mathbf{w}^T \mathbf{x}^{(i)} > 0$
- ▶ $\forall i, \mathbf{x}^{(i)} \in C_2 \Rightarrow \mathbf{w}^T \mathbf{x}^{(i)} < 0$
- ▶ $f(\mathbf{x}; \mathbf{w}) = \text{sign}(\mathbf{w}^T \mathbf{x})$

Perceptron criterion

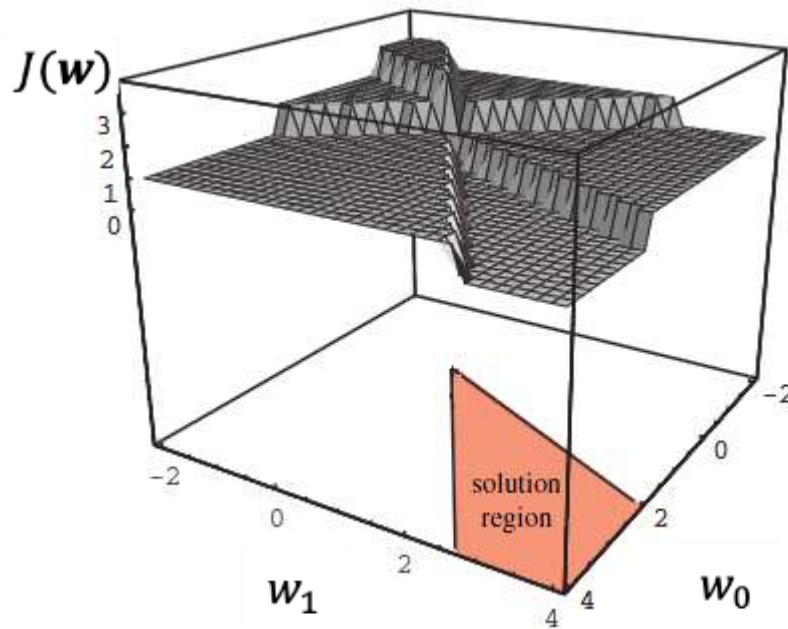
?

$$J_P(\mathbf{w}) = - \sum_{i \in \mathcal{M}} \mathbf{w}^T \mathbf{x}^{(i)} y^{(i)}$$

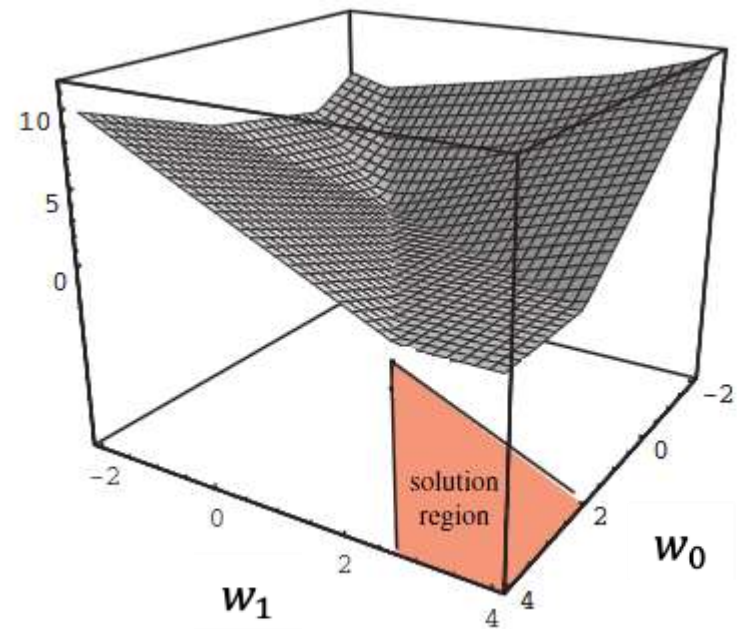
$\mathcal{M}$: subset of training data that are misclassified

Many solutions? Which solution among them?

Cost function



of misclassifications
as a cost function



Perceptron's
cost function

There may be many solutions in these cost functions

[Duda, Hart, and Stork, 2002]

Batch Perceptron

“Gradient Descent” to solve the optimization problem:

$$\mathbf{w}^{t+1} = \mathbf{w}^t - \eta \nabla_{\mathbf{w}} J_P(\mathbf{w}^t)$$

$$\nabla_{\mathbf{w}} J_P(\mathbf{w}) = - \sum_{i \in \mathcal{M}} \mathbf{x}^{(i)} y^{(i)}$$

Batch Perceptron converges in finite number of steps for linearly separable data:

Initialize $\mathbf{w}$

Repeat

$$\mathbf{w} = \mathbf{w} + \eta \sum_{i \in \mathcal{M}} \mathbf{x}^{(i)} y^{(i)}$$

Until $\eta \sum_{i \in \mathcal{M}} \mathbf{x}^{(i)} y^{(i)} < \theta$

Stochastic gradient descent for Perceptron

Single-sample perceptron:

- ▶ If $\mathbf{x}^{(i)}$ is misclassified:

$$\mathbf{w}^{t+1} = \mathbf{w}^t + \eta \mathbf{x}^{(i)} y^{(i)}$$

- ▶ Perceptron convergence theorem: for linearly separable data
 - ▶ If training data are linearly separable, the single-sample perceptron is also guaranteed to find a solution in a finite number of steps

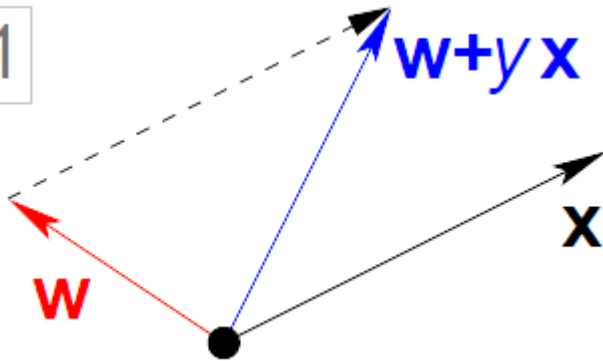
Fixed-Increment single sample Perceptron

```
w ← 0
t ← 0
repeat
  t ← t + 1
  i ← t mod N
  if  $\mathbf{x}^{(i)}$  is misclassified then
     $\mathbf{w} = \mathbf{w} + \mathbf{x}^{(i)} y^{(i)}$ 
Until all patterns properly classified
```

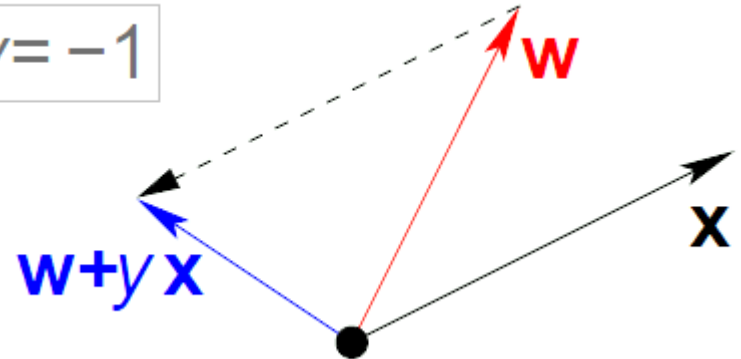
η can be set to 1 and
proof still works →

Example

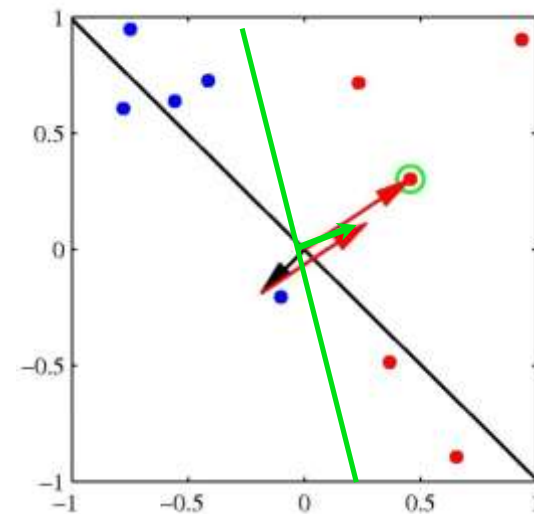
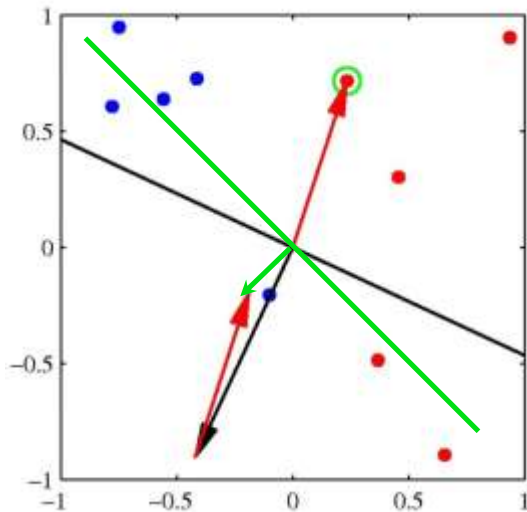
$$y = +1$$



$$y = -1$$



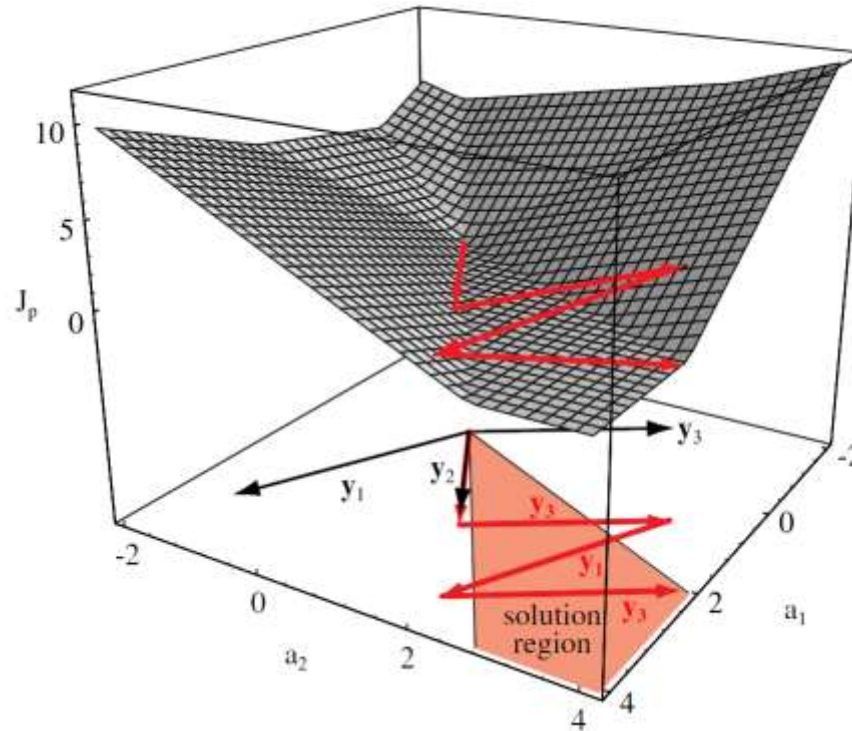
Perceptron: Example



Change w in a direction
that corrects the error

[Bishop]

Convergence of Perceptron



[Duda, Hart & Stork, 2002]

- ? For data sets that are not linearly separable, the single-sample perceptron learning algorithm will never converge

Pocket algorithm

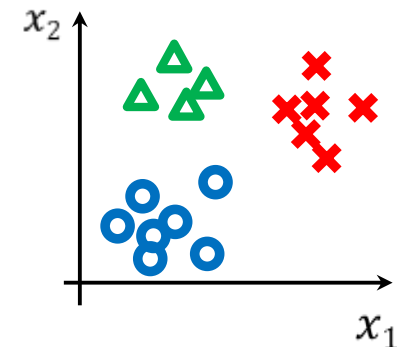
- For the data that are not linearly separable due to noise:
 - Keeps in its pocket the best $\mathbf{w}$ encountered up to now.

```
Initialize  $\mathbf{w}$ 
for  $t = 1, \dots, T$ 
     $i \leftarrow t \bmod N$ 
    if  $\mathbf{x}^{(i)}$  is misclassified then
         $\mathbf{w}^{new} = \mathbf{w} + \mathbf{x}^{(i)} y^{(i)}$ 
        if  $E_{train}(\mathbf{w}^{new}) < E_{train}(\mathbf{w})$  then
             $\mathbf{w} = \mathbf{w}^{new}$ 
end
```

$$E_{train}(\mathbf{w}) = \frac{1}{N} \sum_{n=1}^N [\text{sign}(\mathbf{w}^T \mathbf{x}^{(n)}) \neq y^{(n)}]$$

Multi-class classification

- ▶ Solutions to multi-category problems:
 - ▶ Converting the problem to a set of two-class problems
 - ▶ Extend the learning algorithm to support multi-class:
 - ▶ A function $f_i(\mathbf{x})$ for each class i is found
 - $\hat{y} = \operatorname{argmax}_{i=1,\dots,c} f_i(\mathbf{x})$



Feed back

? <https://forms.gle/vKRbyVVsWRKcZuqr8>



Multi-class classification

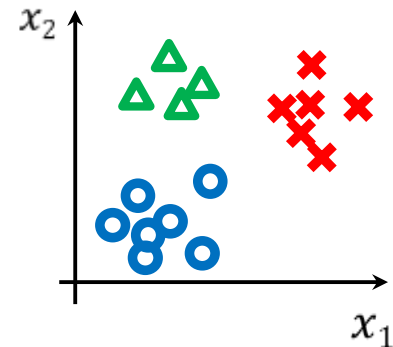
► Solutions to multi-category problems:

- Converting the problem to a set of two-class problems

- Extend the learning algorithm to support multi-class:

- A function $f_i(x)$ for each class i is found

- $\hat{y} = \operatorname{argmax}_{i=1,\dots,c} f_i(x)$



Converting multi-class problem to a set of two-class problems

▶ “one versus rest” or “one against all”

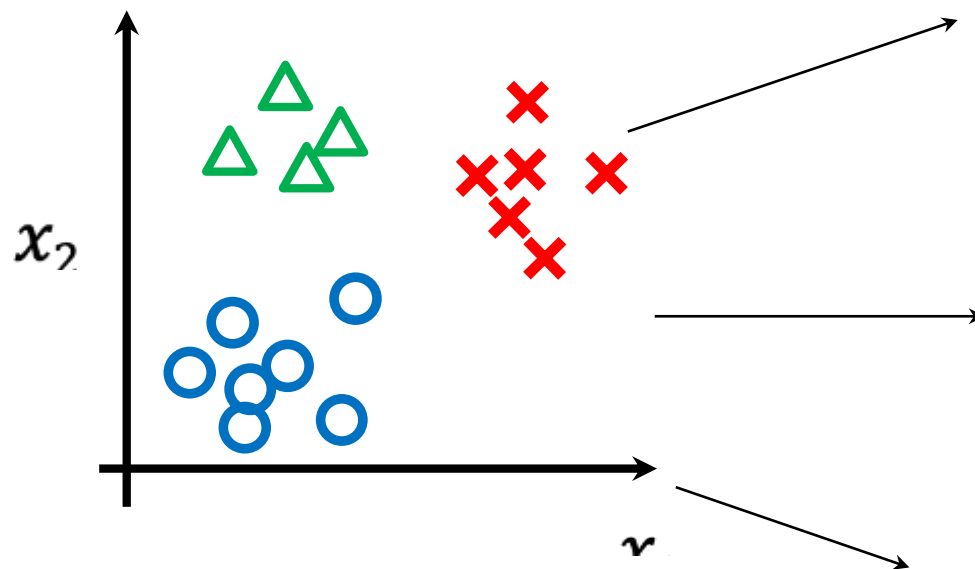
- ▶ For each class C_i , a linear discriminant function that separates samples of C_i from all the other samples is found.
 - ▶ Totally linearly separable

▶ “one versus one”

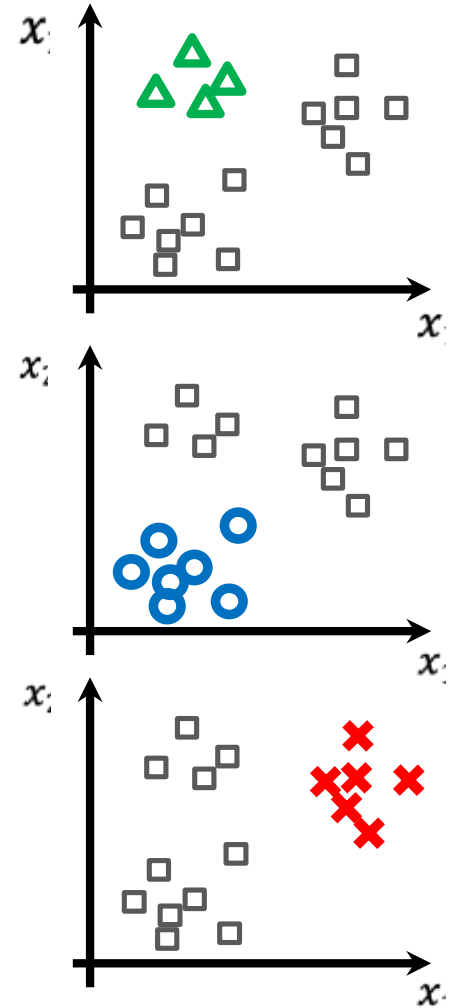
- ▶ $c(c - 1)/2$ linear discriminant functions are used, one to separate samples of a pair of classes.
 - ▶ Pairwise linearly separable

Multi-class classification

One-vs-all (one-vs-rest)

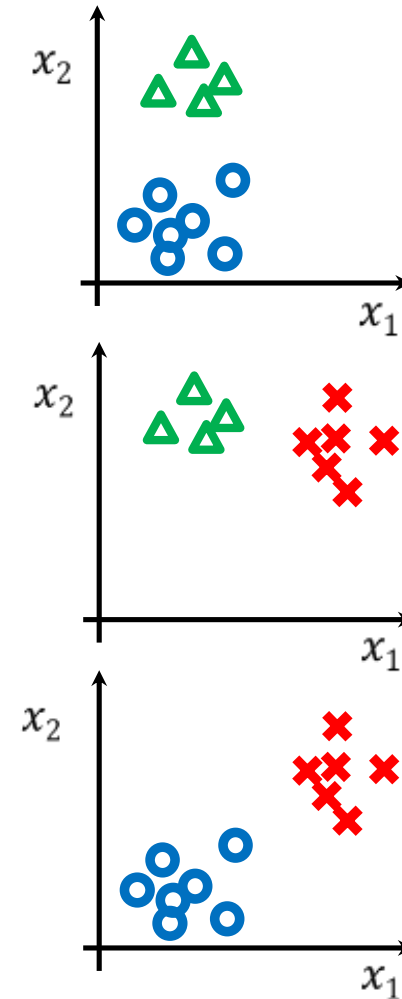
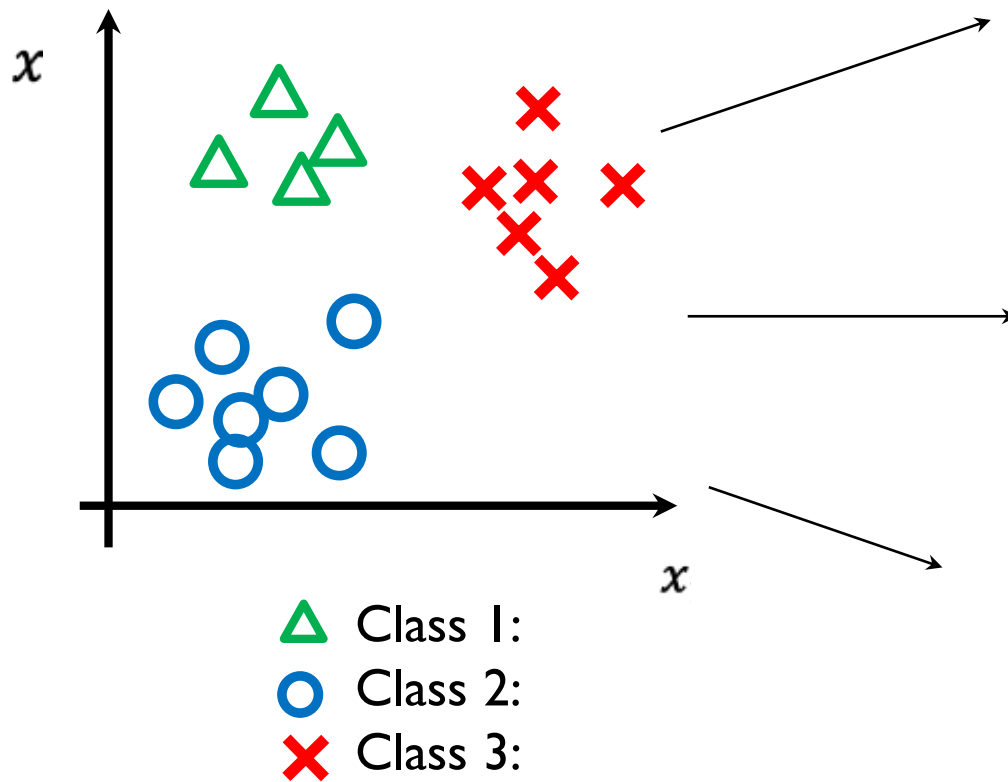


- △ Class 1:
- Class 2:
- × Class 3:



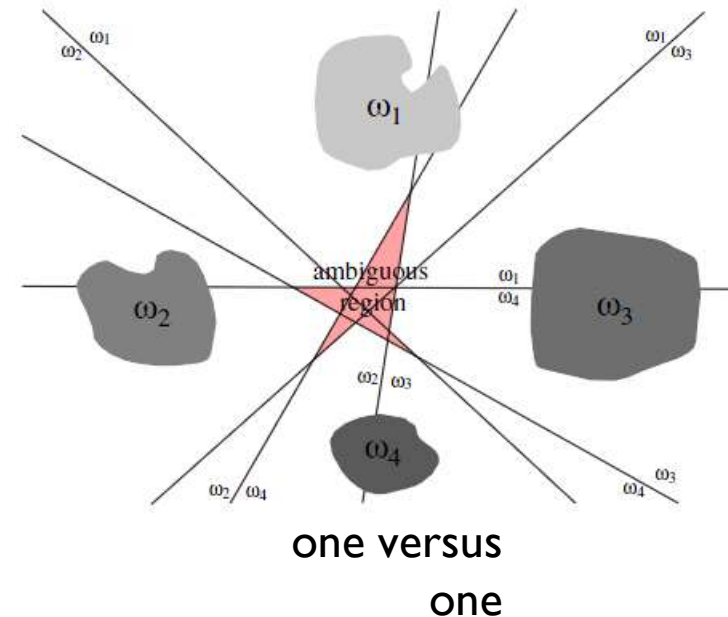
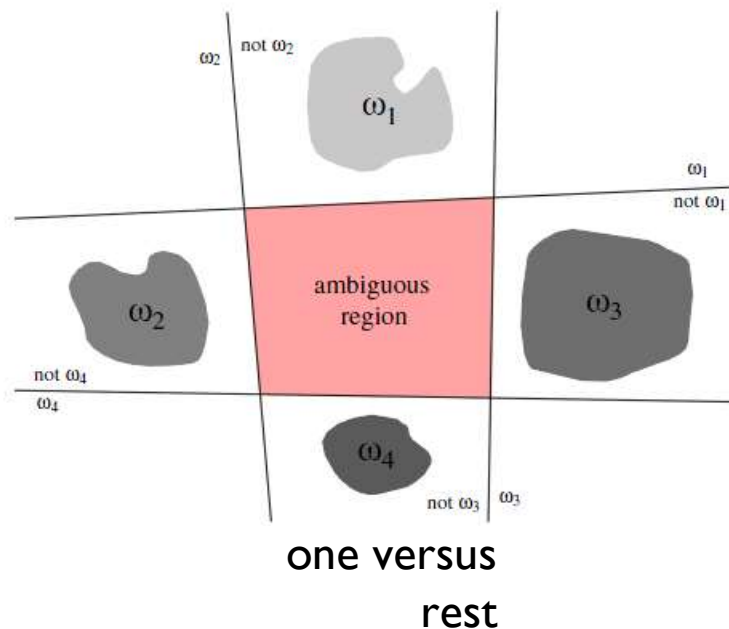
Multi-class classification

One-vs-one



Multi-class classification: ambiguity

Converting the multi-class problem to a set of two-class problems can lead to **regions in which the classification is undefined**



[Duda, Hart & Stork, 2002]

Multi-class classification

► Solutions to multi-category problems:

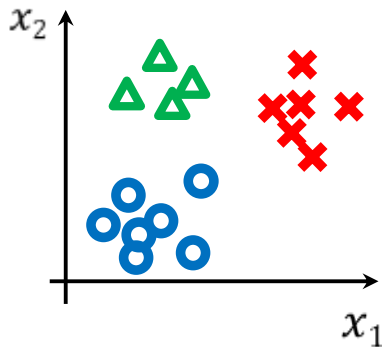
► Converting the problem to a set of two-class problems

► Extend the learning algorithm to support multi-class:

► A function $f_i(\mathbf{x})$ for each class i is found

$$\square \hat{y} = \operatorname{argmax}_{i=1,\dots,c} f_i(\mathbf{x})$$

$\mathbf{x}$ is assigned to class C_i if $f_i(\mathbf{x}) > f_j(\mathbf{x}) \quad \forall j \neq i$



Discriminant Functions

- ▶ **Discriminant functions:** A discriminant function $f_i(\mathbf{x})$ for each class $\mathcal{C}_i$ ($i = 1, \dots, K$):

- ▶ $\mathbf{x}$ is assigned to class $\mathcal{C}_i$ if:

$$f_i(\mathbf{x}) > f_j(\mathbf{x}) \quad \forall j \neq i$$

- ▶ Thus, we can easily divide the feature space into K decision regions

$$\forall \mathbf{x}, f_i(\mathbf{x}) > f_j(\mathbf{x}) \quad \forall j \neq i \Rightarrow \mathbf{x} \in \mathcal{R}_i$$

$\mathcal{R}_i$: Region of the i -th class

- ▶ **Decision surfaces (or boundaries)** can also be found using discriminant functions

- ▶ Boundary of the $\mathcal{R}_i$ and $\mathcal{R}_j$ separating samples of these two categories:

$$\forall \mathbf{x}, f_i(\mathbf{x}) = f_j(\mathbf{x})$$

Discriminant functions

- ▶ **Discriminant function** can directly assign each vector x to a specific class k
- ▶ A popular way of representing a classifier
 - ▶ Many classification methods are based on discriminant functions
- ▶ Assumption: the classes are taken to be disjoint
 - ▶ The input space is thereby divided into **decision regions**
 - ▶ boundaries are called **decision boundaries** or decision surfaces.

Discriminant Functions: Two-Category

?

- ▶ Decision surface: $f(\mathbf{x}) = 0$
- ▶ For two-category problem, we can only find a function $f : \mathbb{R}^d \rightarrow \mathbb{R}$
 - ▶ $f_1(\mathbf{x}) = f(\mathbf{x})$
 - ▶ $f_2(\mathbf{x}) = -f(\mathbf{x})$

Multi-class classification: linear machine

- ▶ A discriminant function $f_i(\mathbf{x}) = \mathbf{w}_i^T \mathbf{x} + w_{i0}$ for each class $\mathcal{C}_i$ ($i = 1, \dots, K$):
 - ▶ $\mathbf{x}$ is assigned to class $\mathcal{C}_i$ if:

$$f_i(\mathbf{x}) > f_j(\mathbf{x}) \quad \forall j \neq i$$

- ▶ **Decision surfaces** (boundaries) can also be found using discriminant functions
 - ▶ Boundary of the contiguous $\mathcal{R}_i$ and $\mathcal{R}_j$: $\forall \mathbf{x}, f_i(\mathbf{x}) = f_j(\mathbf{x})$
 - ▶ $(\mathbf{w}_i - \mathbf{w}_j)^T \mathbf{x} + (w_{i0} - w_{j0}) = 0$
- ▶ Decision regions are convex

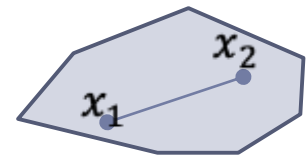
Multi-class classification: linear machine

Decision regions are convex

- Linear machines are most suitable for problems where $p(x|\mathcal{C}_i)$ are unimodal.

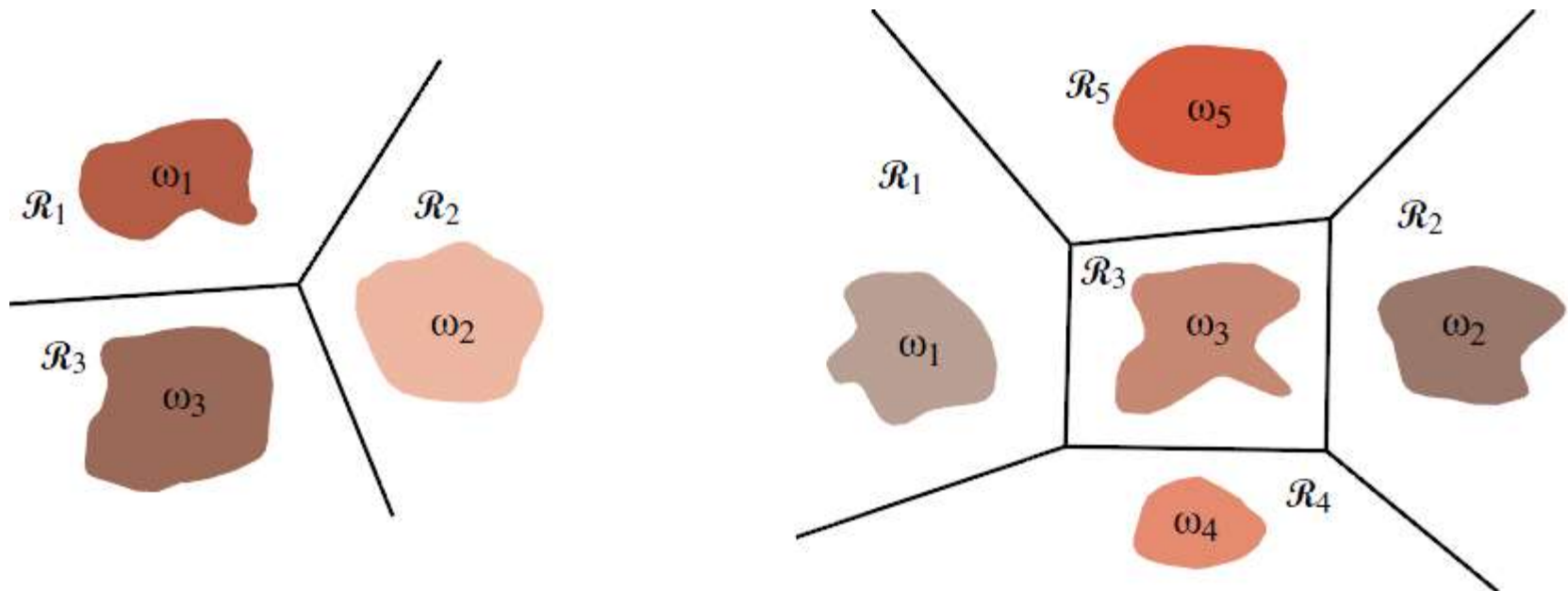
$$x_1, x_2 \in \mathcal{R}_i \Rightarrow \forall j \neq i, f_i(x_1) \geq f_j(x_1) \\ f_i(x_2) \geq f_j(x_2)$$

$$\Rightarrow \alpha f_i(x_1) + (1 - \alpha) f_i(x_2) \geq \alpha f_j(x_1) + (1 - \alpha) f_j(x_2) \\ f_i \text{ is linear} \Rightarrow f_i(\alpha x_1 + (1 - \alpha)x_2) \geq f_j(\alpha x_1 + (1 - \alpha)x_2) \\ \Rightarrow \alpha x_1 + (1 - \alpha)x_2 \in \mathcal{R}_i$$



Convex region definition: $\forall x_1, x_2 \in \mathcal{R}, 0 \leq \alpha \leq 1 \Rightarrow \alpha x_1 + (1 - \alpha)x_2 \in \mathcal{R}$

Multi-class classification: linear machine



[Duda, Hart & Stork, 2002]

Perceptron: multi-class

?

$$\hat{y} = \operatorname{argmax}_{i=1,\dots,c} \mathbf{w}_i^T \mathbf{x}$$

$$J_P(\mathbf{W}) = - \sum_{i \in \mathcal{M}} \left(\mathbf{w}_{y^{(i)}} - \mathbf{w}_{\hat{y}^{(i)}} \right)^T \mathbf{x}^{(i)}$$

$\mathcal{M}$: subset of training data that are misclassified

$$\mathcal{M} = \{i | \hat{y}^{(i)} \neq y^{(i)}\}$$

Initialize $\mathbf{W} = [\mathbf{w}_1, \dots, \mathbf{w}_c]$, $k \leftarrow 0$

repeat

$k \leftarrow (k + 1) \bmod N$

if $\mathbf{x}^{(i)}$ is misclassified then

$$\mathbf{w}_{\hat{y}^{(i)}} = \mathbf{w}_{\hat{y}^{(i)}} - \mathbf{x}^{(i)}$$

$$\mathbf{w}_{y^{(i)}} = \mathbf{w}_{y^{(i)}} + \mathbf{x}^{(i)}$$

Until all patterns properly classified

Linear Discriminant Analysis (LDA)

► Fisher's Linear Discriminant Analysis :

► Dimensionality reduction

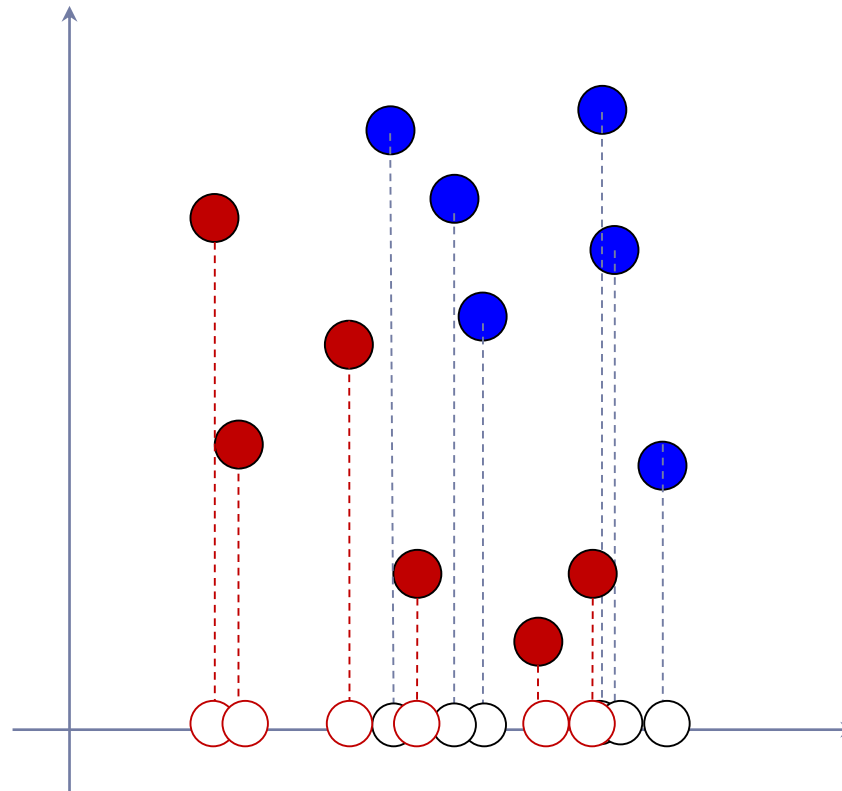
- Finds linear combinations of features with large ratios of between-groups scatters to within-groups scatters (as discriminant new variables)

► Classification

- Predicts the class of an observation x by first projecting it to the space of discriminant variables and then classifying it in this space

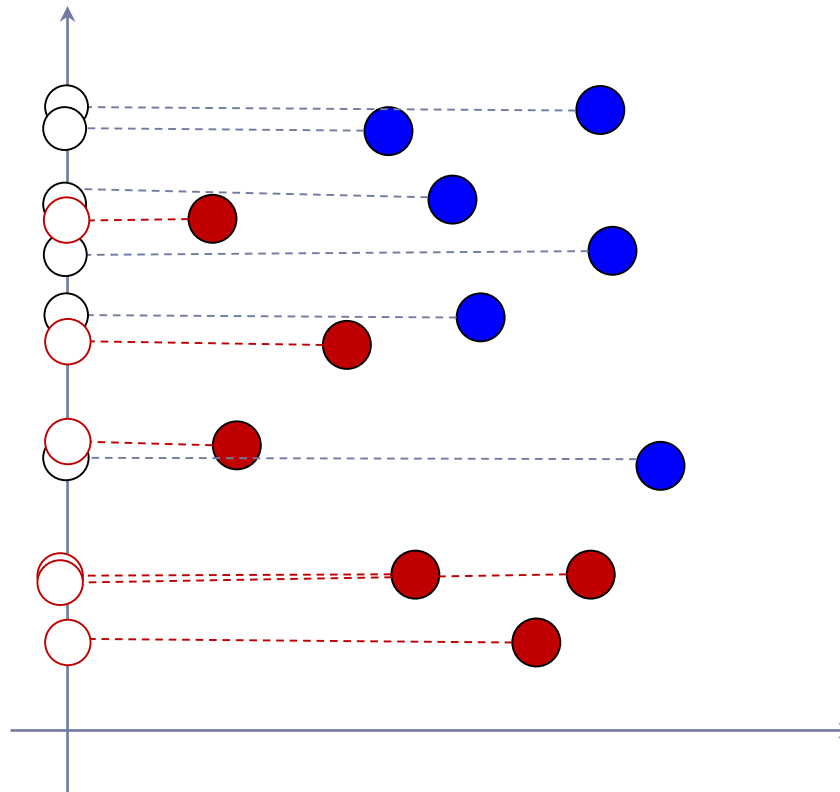
Good Projection for Classification

- ? What is a good criterion?
- ? Separating different classes in the projected space



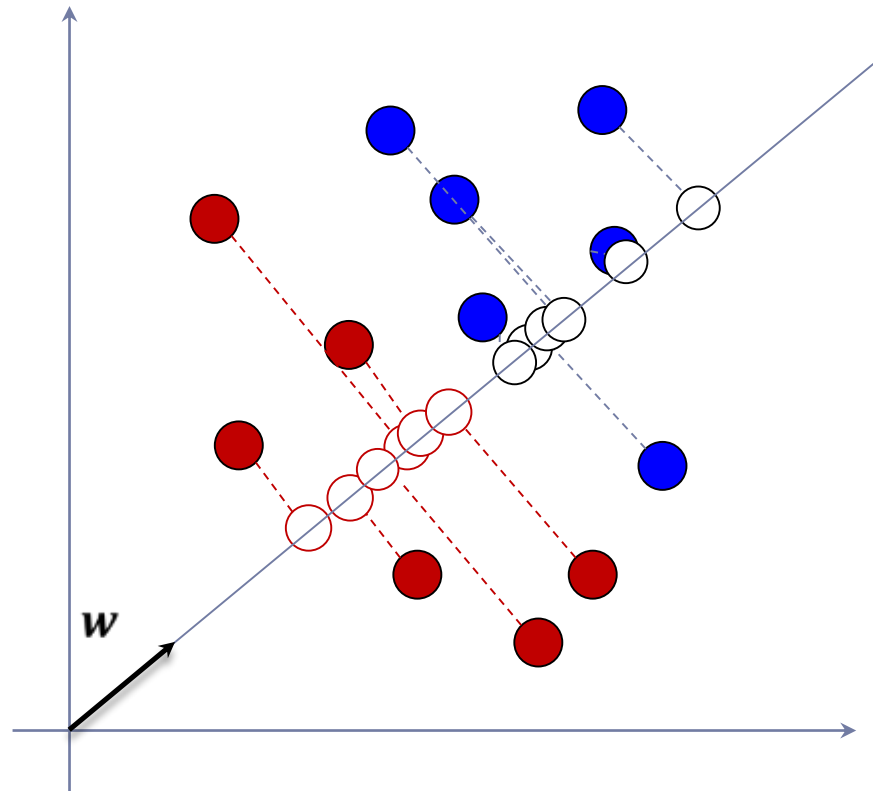
Good Projection for Classification

- ? What is a good criterion?
- ? Separating different classes in the projected space



Good Projection for Classification

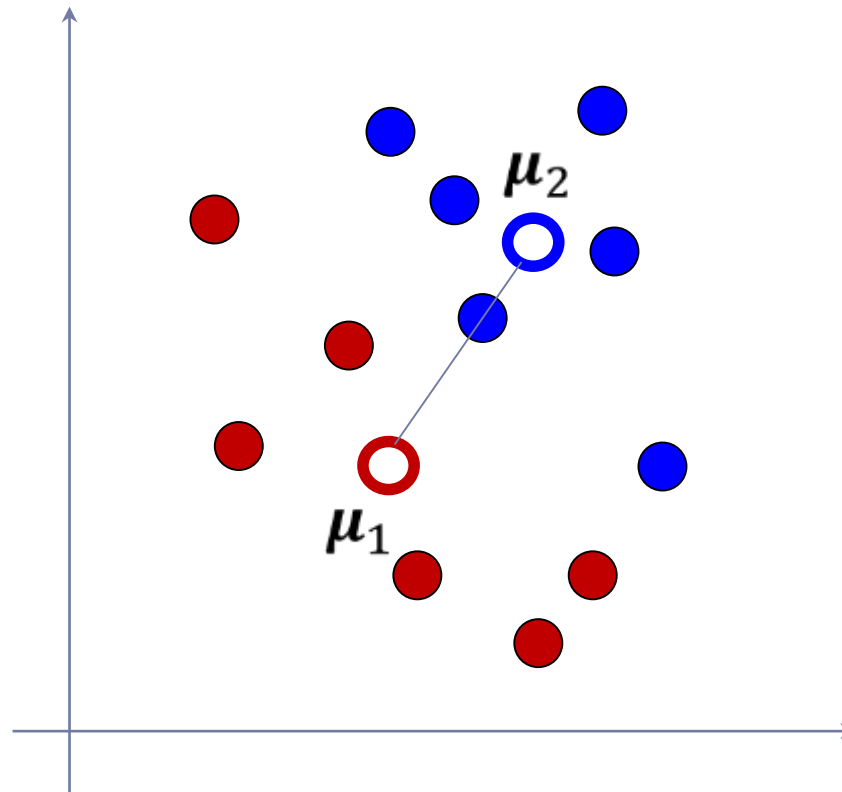
- ? What is a good criterion?
- ? Separating different classes in the projected space



Good Projection for Classification

? Between-class distance

? Squared distance between the centroids of different classes



LDA Problem

Problem definition:

- ▶ $K = 2$ classes
 - ▶ $\{(\mathbf{x}^{(i)}, y^{(i)})\}_{i=1}^N$ training samples with N_1 samples from the first class ($\mathcal{C}_1$) and N_2 samples from the second class ($\mathcal{C}_2$)
 - ▶ Goal: finding the best direction $\mathbf{w}$ that we hope to enable accurate classification
-
- ▶ The projection of sample $\mathbf{x}$ onto a line in direction $\mathbf{w}$ is $\mathbf{w}^T \mathbf{x}$
 - ▶ What is the measure of the separation between the projected points of different classes?

Measure of Separation in the Projected Direction

▶ The direction of the line jointing the class means is the solution of the following problem:

▶ Maximizes the separation of the projected class means

$$\begin{aligned} \max_{\mathbf{w}} J(\mathbf{w}) &= (\mu'_1 - \mu'_2)^2 \\ \text{s. t. } \|\mathbf{w}\| &= 1 \end{aligned}$$

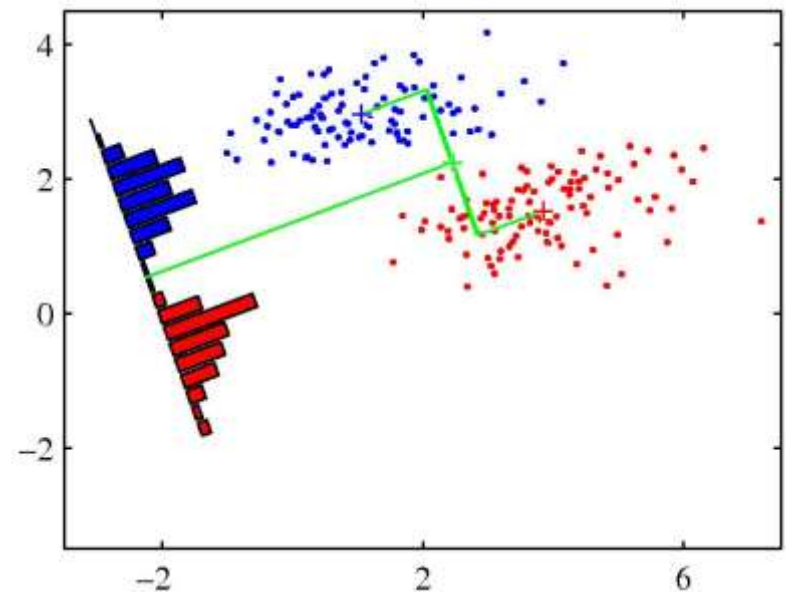
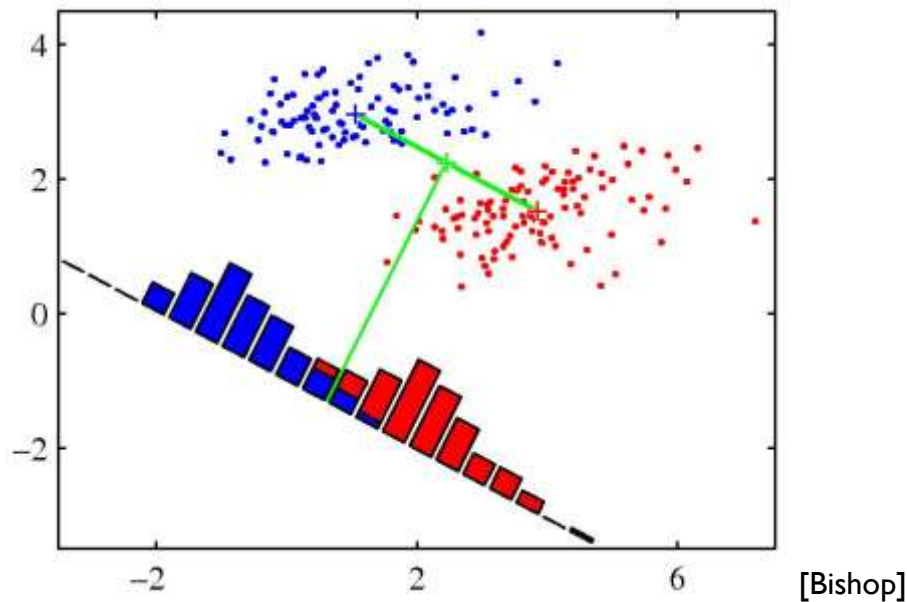
$$\begin{aligned} \mu'_1 &= \mathbf{w}^T \boldsymbol{\mu}_1 & \boldsymbol{\mu}_1 &= \frac{\sum_{\mathbf{x}^{(i)} \in \mathcal{C}_1} \mathbf{x}^{(i)}}{N_1} \\ \mu'_2 &= \mathbf{w}^T \boldsymbol{\mu}_2 & \boldsymbol{\mu}_2 &= \frac{\sum_{\mathbf{x}^{(i)} \in \mathcal{C}_2} \mathbf{x}^{(i)}}{N_2} \end{aligned}$$

▶ What is the problem with the criteria considering only $|\mu'_1 - \mu'_2|$?

▶ It does not consider the variances of the classes in the projected direction

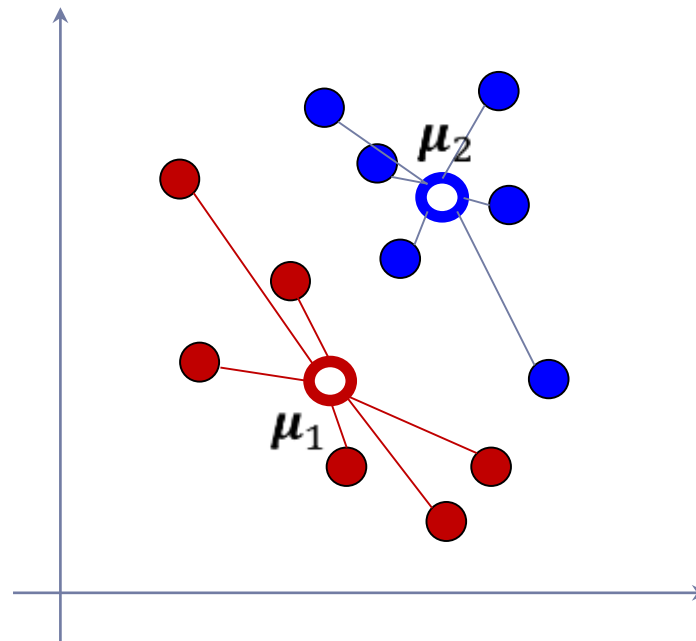
Measure of Separation in the Projected Direction

- ▶ Is the direction of the line joining the class means a good candidate for w ?



Good Projection for Classification

- ? Between-class distance
 - ? Squared distance between the centroids/means of different classes
- ? Within-class distance
 - ? Accumulated squared distance of an instance to the centroid/ mean of its class



LDA Criteria

- ▶ Fisher idea: maximize a function that will give
 - ▶ large separation between the projected class means
 - ▶ while also achieving a small variance within each class, thereby minimizing the class overlap.

$$J(\mathbf{w}) = \frac{|\mu'_1 - \mu'_2|^2}{s_1'^2 + s_2'^2}$$

LDA Criteria

• The scatters of projected data are:

$$s_1'^2 = \sum_{\mathbf{x}^{(i)} \in \mathcal{C}_1} (\mathbf{w}^T \mathbf{x}^{(i)} - \mathbf{w}^T \boldsymbol{\mu}_1)^2$$

$$s_2'^2 = \sum_{\mathbf{x}^{(i)} \in \mathcal{C}_2} (\mathbf{w}^T \mathbf{x}^{(i)} - \mathbf{w}^T \boldsymbol{\mu}_1)^2$$

LDA Criteria

$$J(\mathbf{w}) = \frac{|\mu'_1 - \mu'_2|^2}{s_1'^2 + s_2'^2}$$

$$\begin{aligned} |\mu'_1 - \mu'_2|^2 &= |\mathbf{w}^T \boldsymbol{\mu}_1 - \mathbf{w}^T \boldsymbol{\mu}_2|^2 \\ &= \mathbf{w}^T (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)(\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)^T \mathbf{w} \end{aligned}$$

$$\begin{aligned} s_1'^2 &= \sum_{\mathbf{x}^{(i)} \in \mathcal{C}_1} (\mathbf{w}^T \mathbf{x}^{(i)} - \mathbf{w}^T \boldsymbol{\mu}_1)^2 \\ &= \mathbf{w}^T \left(\sum_{\mathbf{x}^{(i)} \in \mathcal{C}_1} (\mathbf{x}^{(i)} - \boldsymbol{\mu}_1)(\mathbf{x}^{(i)} - \boldsymbol{\mu}_1)^T \right) \mathbf{w} \end{aligned}$$

LDA Criteria

$$J(\mathbf{w}) = \frac{\mathbf{w}^T \mathbf{S}_B \mathbf{w}}{\mathbf{w}^T \mathbf{S}_W \mathbf{w}}$$

Between-class
scatter matrix

$$\mathbf{S}_B = (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)(\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)^T$$

Within-class
scatter matrix

$$\mathbf{S}_W = \mathbf{S}_1 + \mathbf{S}_2$$

$$\mathbf{S}_1 = \sum_{\mathbf{x}^{(i)} \in \mathcal{C}_1} (\mathbf{x}^{(i)} - \boldsymbol{\mu}_1)(\mathbf{x}^{(i)} - \boldsymbol{\mu}_1)^T$$

$$\mathbf{S}_2 = \sum_{\mathbf{x}^{(i)} \in \mathcal{C}_2} (\mathbf{x}^{(i)} - \boldsymbol{\mu}_2)(\mathbf{x}^{(i)} - \boldsymbol{\mu}_2)^T$$

scatter matrix = $N \times$ covariance matrix

LDA Derivation

$$J(\mathbf{w}) = \frac{\mathbf{w}^T \mathbf{S}_B \mathbf{w}}{\mathbf{w}^T \mathbf{S}_W \mathbf{w}}$$

$$\frac{\partial J(\mathbf{w})}{\partial \mathbf{w}} = \frac{\frac{\partial \mathbf{w}^T \mathbf{S}_B \mathbf{w}}{\partial \mathbf{w}} \times \mathbf{w}^T \mathbf{S}_W \mathbf{w} - \frac{\partial \mathbf{w}^T \mathbf{S}_W \mathbf{w}}{\partial \mathbf{w}} \times \mathbf{w}^T \mathbf{S}_B \mathbf{w}}{(\mathbf{w}^T \mathbf{S}_W \mathbf{w})^2} = \frac{(2\mathbf{S}_B \mathbf{w}) \mathbf{w}^T \mathbf{S}_W \mathbf{w} - (2\mathbf{S}_W \mathbf{w}) \mathbf{w}^T \mathbf{S}_B \mathbf{w}}{(\mathbf{w}^T \mathbf{S}_W \mathbf{w})^2}$$

$$\frac{\partial J(\mathbf{w})}{\partial \mathbf{w}} = 0 \Rightarrow \mathbf{S}_B \mathbf{w} = \lambda \mathbf{S}_W \mathbf{w}$$

LDA Derivation

$$? \mathbf{S}_B \mathbf{w} = \lambda \mathbf{S}_W \mathbf{w} \xrightarrow{\text{If } \mathbf{S}_W \text{ is full-rank}} \mathbf{S}_W^{-1} \mathbf{S}_B \mathbf{w} = \lambda \mathbf{w}$$

$\mathbf{S}_B \mathbf{w}$ (for any vector $\mathbf{w}$) points in the same direction as $\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2$:

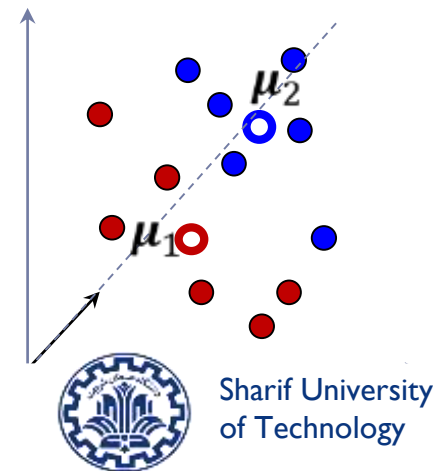
$$\mathbf{S}_B \mathbf{w} = (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)(\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)^T \mathbf{w} \propto (\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)$$

$$\mathbf{w} \propto \mathbf{S}_W^{-1}(\boldsymbol{\mu}_1 - \boldsymbol{\mu}_2)$$

Thus, we can solve the eigenvalue problem immediately

LDA Algorithm

- Find μ_1 and μ_2 as the mean of class 1 and 2 respectively
- Find S_1 and S_2 as scatter matrix of class 1 and 2 respectively
- $S_W = S_1 + S_2$
- $S_B = (\mu_1 - \mu_2)(\mu_1 - \mu_2)^T$
- Feature Extraction
 - $w = S_W^{-1}(\mu_1 - \mu_2)$ as the eigenvector corresponding to the largest eigenvalue of $S_W^{-1}S_B$
- Classification
 - $w = S_W^{-1}(\mu_1 - \mu_2)$
 - Using a threshold on $w^T x$, we can classify x



Multi-class classification: target coding scheme

Target values:

- ▶ Binary classification: a target variable $y \in \{0,1\}$
- ▶ Multiple classes ($K > 2$):
 - ▶ Target class $\mathcal{C}_j$: $y_j = 1$
 $\forall i \neq j \ y_i = 0$ } One of K coding

Resources

?

$$J(W) = \sum_{k=1}^K \|Xw_k - y_k\|_2^2$$

SE cost function for multi-class

$$X = \begin{bmatrix} 1 & x_1^{(1)} & \dots & x_d^{(1)} \\ 1 & x_1^{(2)} & \dots & x_d^{(2)} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_1^{(n)} & \dots & x_d^{(n)} \end{bmatrix} \quad y_k = \begin{bmatrix} y_k^{(1)} \\ \vdots \\ y_k^{(n)} \end{bmatrix}$$

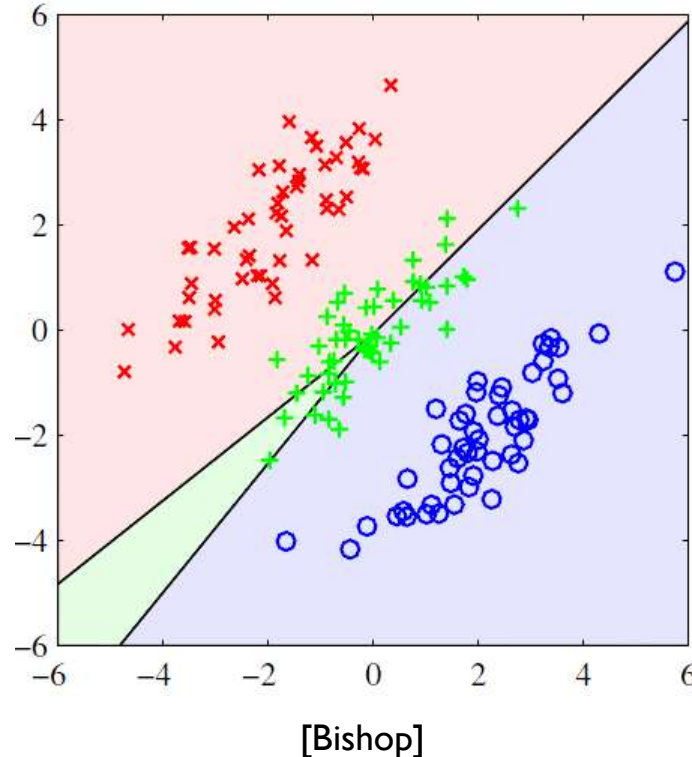
$$W = [w_1 \quad \dots \quad w_K]$$

$$Y = [y_1 \quad \dots \quad y_K]$$

$$\nabla_W J(W) = 0 \Rightarrow \widehat{W} = (X^T X)^{-1} X Y$$

SSE for multi-class: example

Low performance of the SSE cost function for the classification problem



Exercise

Think about multi-class perceptron, LDA



Multi-class LDA

▶ $\mathbf{W} = [\mathbf{w}_1 \ \mathbf{w}_2 \ \dots \ \mathbf{w}_{K-1}]$

▶ $\mathbf{x}' = \mathbf{W}^T \mathbf{x}$

▶ Means and scatters after transform $\mathbf{x}' = \mathbf{W}^T \mathbf{x}$:

▶ $\mathbf{S}'_B = \mathbf{W}^T \mathbf{S}_B \mathbf{W}$

▶ $\mathbf{S}'_W = \mathbf{W}^T \mathbf{S}_W \mathbf{W}$

Multi-Class LDA: Objective Function

- ▶ We seek a transformation matrix W that in some sense “maximizes the ratio of the between-class scatter to the within-class scatter”.
- ▶ A simple scalar measure of scatter is the **determinant of the scatter matrix**.

Multi-class LDA (MDA)

- ▶ $C > 2$: the natural generalization of LDA involves $C - 1$ discriminant functions.
- ▶ The projection from a d -dimensional space to a $(C - 1)$ -dimensional space (tacitly assumed that $d \geq C$).

$$\mathbf{S}_W = \sum_{j=1}^K \mathbf{S}_j$$

$$\mathbf{S}_B = \sum_{j=1}^K N_j (\boldsymbol{\mu}_j - \boldsymbol{\mu})(\boldsymbol{\mu}_j - \boldsymbol{\mu})^T$$

$$\boldsymbol{\mu}_j = \frac{\sum_{\mathbf{x}^{(i)} \in \mathcal{C}_j} \mathbf{x}^{(i)}}{N_j} \quad j = 1, \dots, K$$

$$\boldsymbol{\mu} = \frac{\sum_{i=1}^N \mathbf{x}^{(i)}}{N}$$

$$\mathbf{S}_j = \sum_{\mathbf{x}^{(i)} \in \mathcal{C}_j} (\mathbf{x}^{(i)} - \boldsymbol{\mu}_j)(\mathbf{x}^{(i)} - \boldsymbol{\mu}_j)^T \quad j = 1, \dots, K$$

Multi-class LDA: Objective function

?

$$J(\mathbf{W}) = \frac{|\mathbf{W}^T \mathbf{S}_B \mathbf{W}|}{|\mathbf{W}^T \mathbf{S}_W \mathbf{W}|} \longrightarrow \text{determinant}$$

- ▶ The solution of the problem where $\mathbf{W} = [\mathbf{w}_1 \ \mathbf{w}_2 \ \dots \ \mathbf{w}_{C-1}]$:
$$\mathbf{S}_B \mathbf{w}_i = \lambda_i \mathbf{S}_W \mathbf{w}_i$$
- ▶ It is a generalized eigenvectors problem.

Objective Function: Trace Ratio

- ▶ maximizing between-class squared distances and minimizing within-class squared distances:

$$J(\mathbf{W}) = \frac{\text{tr}(\mathbf{W}^T \mathbf{S}_B \mathbf{W})}{\text{tr}(\mathbf{W}^T \mathbf{S}_W \mathbf{W})}$$

- ▶ Between-class distance = trace of between-class scatter
Within-class distance = trace of within-class scatter

Objective Function: Trace Ratio

$$\mathbf{S}_W = \sum_{j=1}^C \sum_{\mathbf{x}^{(i)} \in \mathcal{C}_j} (\mathbf{x}^{(i)} - \boldsymbol{\mu}_j)(\mathbf{x}^{(i)} - \boldsymbol{\mu}_j)^T$$

$$\mathbf{S}_B = \sum_{j=1}^C N_j (\boldsymbol{\mu}_j - \boldsymbol{\mu})(\boldsymbol{\mu}_j - \boldsymbol{\mu})^T$$

$$\text{tr}(\mathbf{S}_W) = \sum_{j=1}^C \sum_{\mathbf{x}^{(i)} \in \mathcal{C}_j} (\mathbf{x}^{(i)} - \boldsymbol{\mu}_j)^T (\mathbf{x}^{(i)} - \boldsymbol{\mu}_j)$$

$$= \sum_{j=1}^C \sum_{\mathbf{x}^{(i)} \in \mathcal{C}_j} \|\mathbf{x}^{(i)} - \boldsymbol{\mu}_j\|^2$$

$$\text{tr}(\mathbf{S}_B) = \sum_{j=1}^C N_j (\boldsymbol{\mu}_j - \boldsymbol{\mu})^T (\boldsymbol{\mu}_j - \boldsymbol{\mu}) = \sum_{j=1}^C N_j \|\boldsymbol{\mu}_j - \boldsymbol{\mu}\|^2$$

Objective Function: Trace Ratio vs. Ratio Trace

$$tr(S_w) = \sum_{i=1}^c \sum_{j \in \text{class } i} \sum_{k \in \text{class } i} \| \mathbf{x}_j - \mathbf{x}_k \|^2$$

$$tr(S_b) = \sum_{i=1}^c \sum_{\substack{j=1 \\ j \neq i}}^c \sum_{j \in \text{class } i} \sum_{k \in \text{class } j} \| \mathbf{x}_j - \mathbf{x}_k \|^2$$

Multi-class LDA: Other objective function

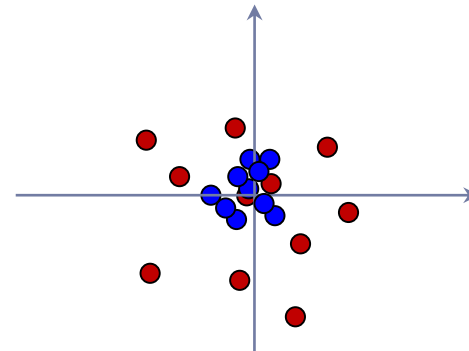
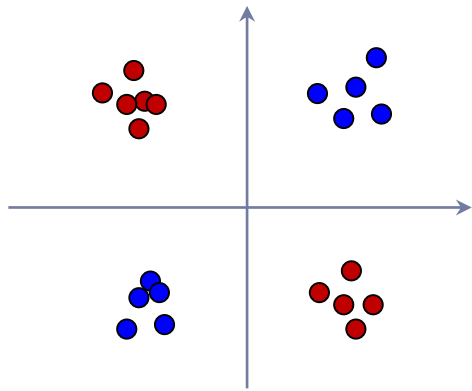
- ▶ There are many possible choices of criterion for multi-class LDA, e.g.:

$$J(W) = \text{tr}(\mathbf{S}_W'^{-1} \mathbf{S}_B') = \text{tr}((\mathbf{W}^T \mathbf{S}_W \mathbf{W})^{-1} (\mathbf{W}^T \mathbf{S}_B \mathbf{W}))$$

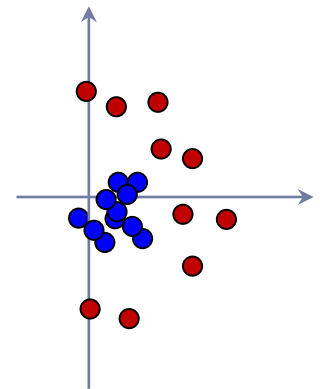
- ▶ The solution is given by solving a generalized eigenvalue problem $\mathbf{S}_W^{-1} \mathbf{S}_b$
 - ▶ Solution: eigen vectors corresponding to the largest eigen values constitute the new variables

LDA Criterion Limitation

- When $\mu_1 = \mu_2$, LDA criterion can not lead to a proper projection ($J(\mathbf{w}) = 0$)
 - However, discriminatory information in the scatter of the data may be helpful



- If classes are non-linearly separable they may have large overlap when projected to any line
- LDA implicitly assumes Gaussian distribution of samples of each class



Issues in LDA

- ▶ Singularity or undersampled problem (when $N < d$)
 - ▶ Example: gene expression data, images, text documents
- ▶ Can reduce dimension only to $d' \leq C - 1$
- ▶ $\text{rank}(\mathbf{S}_B) \leq C - 1$
 - ▶ $\mathbf{S}_B$ is the sum of C matrices $(\boldsymbol{\mu}_j - \boldsymbol{\mu})(\boldsymbol{\mu}_j - \boldsymbol{\mu})^T$ of rank (at most) one and only $C - 1$ of these are independent,
 - ▶ $\Rightarrow$ at most $C - 1$ nonzero eigenvalues and the desired weight vectors correspond to these nonzero eigenvalues.

Summery

- ❓ Although LDA often provide more suitable features for classification tasks, LDA may fails in some situations such as:
 - ❑ when the number of samples per class is small (overfitting problem of LDA)
 - ❑ when the training data non-uniformly sample the underlying distribution
 - ❑ when the number of the desired features is more than $C - 1$
- ❑ Advances in the recent decade:
 - ❑ Semi-supervised feature extraction
 - ❑ Nonlinear dimensionality reduction

Applications

- ? Face recognition
 - ? Belhumeur et al., PAMI'97
- ? Image retrieval
 - ? Swets and Weng, PAMI'96
- ? Gene expression data analysis
 - ? Dudoit et al., JASA'02; Ye et al., TCBB'04
- ? Protein expression data analysis
 - ? Lilien et al., Comp. Bio.'03
- ? Text mining
 - ? Park et al., SIMAX'03; Ye et al., PAMI'04
- ? Medical image analysis
 - ? Dundar, SDM'05



Resources

- ? C. Bishop, “Pattern Recognition and Machine Learning”, Chapter 4.1.
- ? Course CE-717, Dr. M.Soleymani

