



国際融合科学論/先端融合科学論

LECTURE 02

Machine Learning I: conventional machine learning

Dr. Suyong Eum



1) Principal Component Analysis (PCA)

- Feature selections
- Dimension reduction

2) Support Vector Machine (SVM)

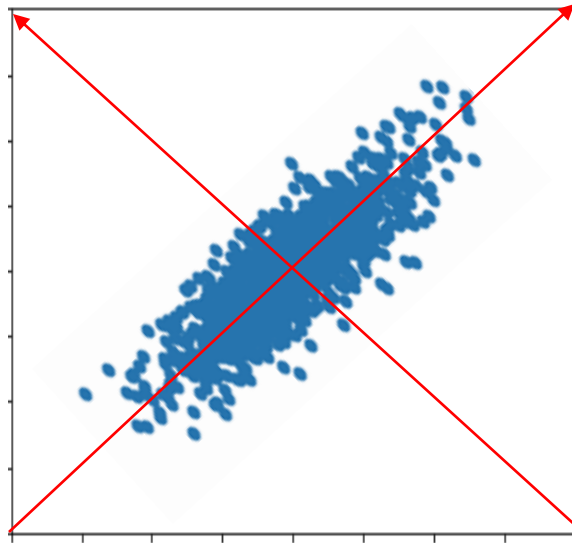
- Hard margin SVM: linear classification
- Kernel trick: nonlinear classification

Principal Component Analysis (PCA)

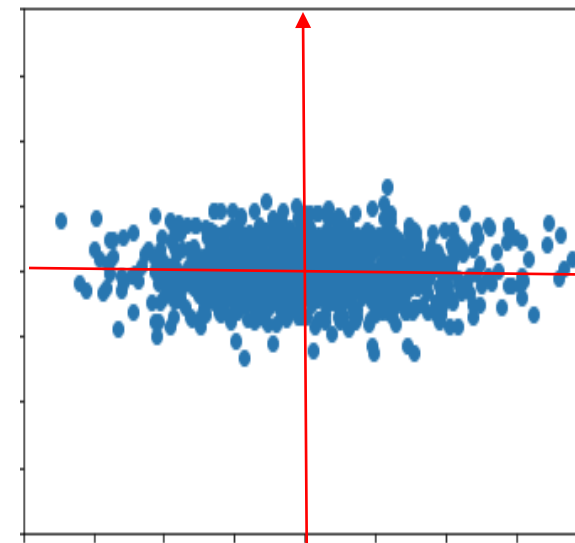
Principal Component Analysis (PCA): definition

A statistical procedure that uses an orthogonal transformation to convert a set of observations of possibly correlated variables into a set of values of linearly uncorrelated variables called principal components.

In Wikipedia:



$$\begin{bmatrix} 8 & 1 \\ 1 & 8 \end{bmatrix}$$

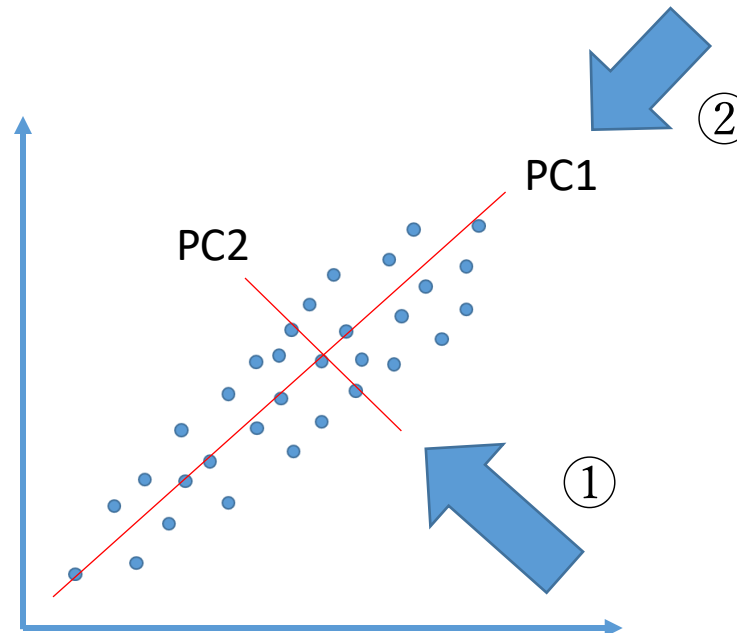


$$\begin{bmatrix} 10 & 0 \\ 0 & 1 \end{bmatrix}$$

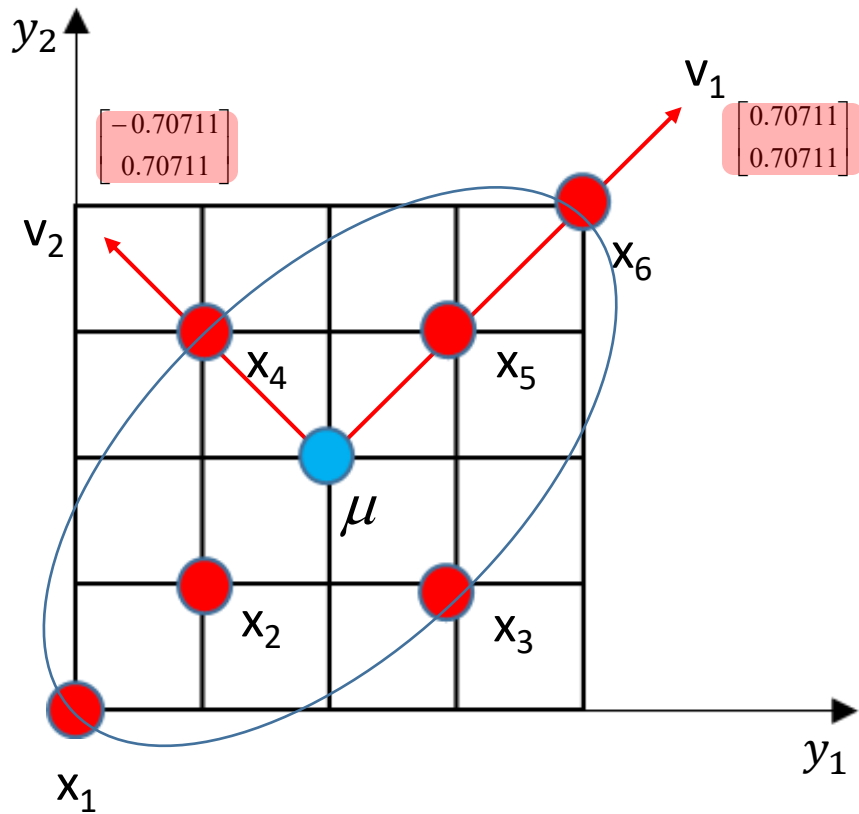
Principal Component Analysis (PCA): intuition

❑ How to select principal components?

- One that captures the largest variance of the data points
- Intuitively speaking, you can observe more data from the direction ① than any other direction, and then from the direction ②, you can observe the data with the least redundancy compared to the direction ①.



How to find the principal components showing the largest variance?



$$X = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \\ x_6 \end{bmatrix} = \begin{bmatrix} -2 & -2 \\ -1 & -1 \\ 1 & -1 \\ -1 & 1 \\ 1 & 1 \\ 2 & 2 \end{bmatrix}$$

Distance to data points from the mean along the axis of " v_1 "
 $= [2\sqrt{2}, \sqrt{2}, 0, 0, \sqrt{2}, 2\sqrt{2}]$ Variance = 4

Distance to data points from the mean along the axis of " v_2 "
 $= [0, 0, \sqrt{2}, \sqrt{2}, 0, 0]$ Variance = 0.8

$\text{cov}(X) = \begin{bmatrix} 2.4 & 1.6 \\ 1.6 & 2.4 \end{bmatrix}$

$\text{cov}(x) = V\Lambda V^T$

```
>> [vec, val] = eig(cov(x))
vec =
    -0.70711    0.70711
     0.70711    0.70711

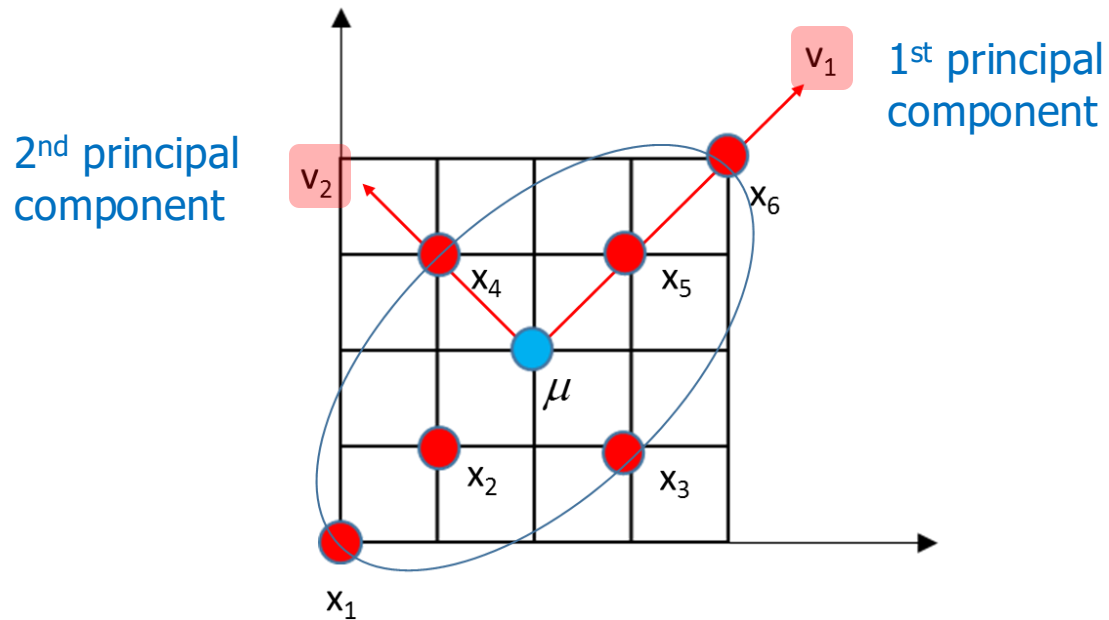
val =
Diagonal Matrix
    0.80000    0
         0    4.00000
```

V

Λ

How to find the principal components showing the largest variance?

- 1) Find the covariance matrix of data points.
- 2) Obtain the eigen values and vectors of the covariance matrix: **eigen value decomposition**.
- 3) Sort the eigen vectors in descending order in terms of their corresponding eigen values.
 - an eigen vector with the largest eigen value becomes the first principal component.



```
>> x
x =

    -2    -2
    -1    -1
     1    -1
    -1     1
     1     1
     2     2

>> cov(x)
ans =

    2.4000    1.6000
    1.6000    2.4000
```

```
>> [vec, val] = eig(cov(x))
vec =

    -0.70711    0.70711
     0.70711    0.70711

val =

    0.80000    0
     0    4.00000

Diagonal Matrix
```

2nd principal component 1st principal component

How to find the principal components showing the largest variance?

- ❑ Actually, there is a more convenient way of doing it, which is called "Singular Value Decomposition" or **SVD**.

Eigen decomposition

$$X^T X = V \Lambda V^T$$

```
>> x
x =
-2 -2
-1 -1
1 -1
-1 1
1 1
2 2

>> cov(x)
ans =
2.4000 1.6000
1.6000 2.4000
```

```
>> [vec, val] = eig(cov(x))
vec =
-0.70711 0.70711
0.70711 0.70711

val =
Diagonal Matrix
0.80000 0
0 4.00000
```

```
>> [vec, val] = eig(transpose(x)*x)
vec =
-0.70711 0.70711
0.70711 0.70711

val =
Diagonal Matrix
4.0000 0
0 20.0000
```

Singular Value Decomposition (SVD)

$$X = U \Sigma V^T$$

$$X^T X = (U \Sigma V^T)^T (U \Sigma V^T)$$

$$= V \Sigma^T U^T U \Sigma V^T$$

$$= V \Sigma^2 V^T$$

$$\Lambda = \Sigma^2$$

$$\Lambda = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{bmatrix}$$

Eigen value

$$\Sigma = \begin{bmatrix} \sqrt{\lambda_1} & & 0 \\ & \ddots & \\ 0 & & \sqrt{\lambda_n} \end{bmatrix}$$

Singular value

How to find the principal components showing the largest variance?

- ❑ Actually, there is a more convenient way of doing it, which is called "Singular Value Decomposition" or **SVD**.

Eigen decomposition

$$X^T X = V \Lambda V^T$$

Singular Value Decomposition (SVD)

$$X = U \Sigma V^T$$

```
>> x
x =
-2 -2
-1 -1
1 -1
-1 1
1 1
2 2

>> cov(x)
ans =
2.4000 1.6000
1.6000 2.4000
```

```
>> [vec, val] = eig(cov(x))
vec =
-0.70711 0.70711
0.70711 0.70711

val =
4.0000 0
0 4.0000

Diagonal Matrix
```

```
>> [vec, val]=eig(transpose(x)*x)
vec =
-0.70711 0.70711
0.70711 0.70711

val =
4.0000 0
0 20.0000

Diagonal Matrix
```

$$\Lambda = \Sigma^2$$

$$4.4721^2 = 20$$

```
>> [U, S, V]=svd(x)
U =
-0.63246 0.00000 0.30819 -0.30819 0.28637 0.57274
-0.31623 -0.00000 -0.63635 0.63635 0.13426 0.26851
0.00000 -0.70711 0.50000 0.50000 0.00000 0.00000
-0.00000 0.70711 0.50000 0.50000 -0.00000 -0.00000
0.31623 0.00000 -0.00399 0.00399 0.94140 -0.11720
0.63246 0.00000 -0.00799 0.00799 -0.11720 0.76560

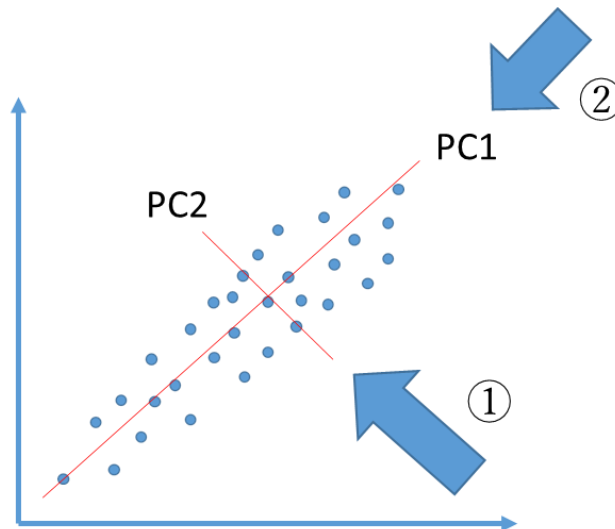
S =
4.4721 0
0 2.0000
0 0
0 0
0 0
0 0

Diagonal Matrix

V =
0.70711 -0.70711
0.70711 0.70711
```

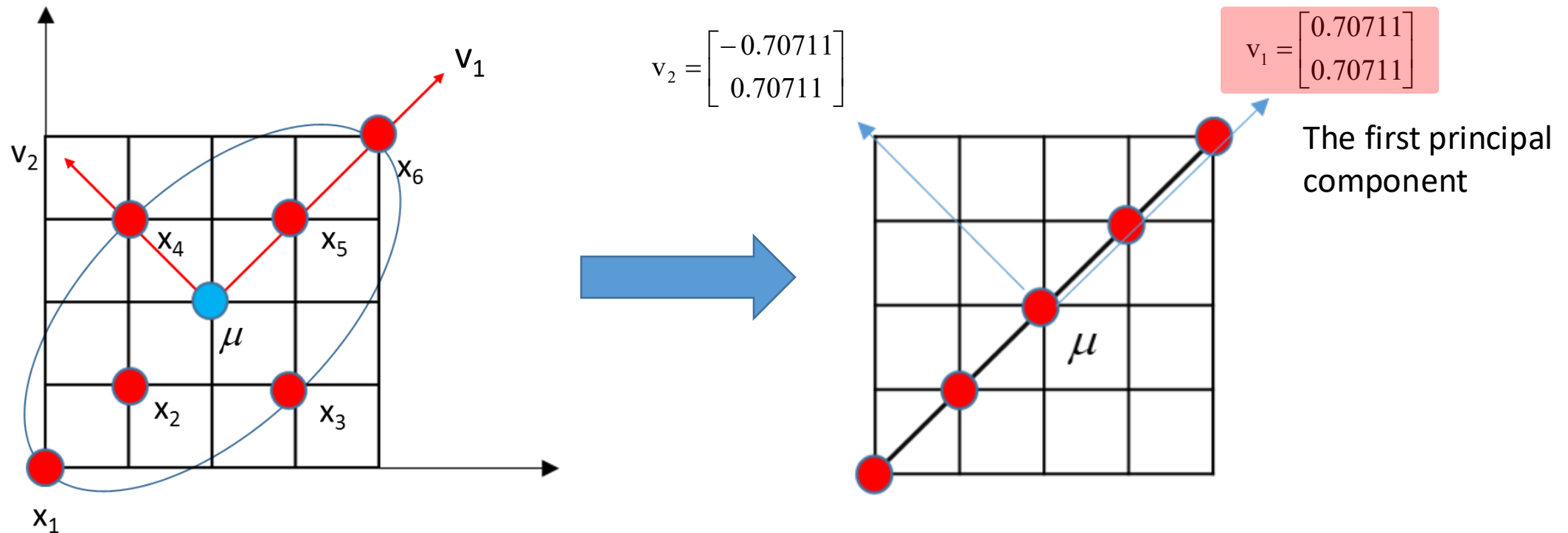
Now we know how to find the principal components

- ❑ Principal Component Analysis (PCA) is nothing but finding principal components of a given data set,
 - Principal components are the directions where you look at the data set, which provides the most information of the data set.
 - They're equivalent to eigen vectors which can be found by SVD or EVD.
 - The eigen value corresponding to each eigen vector represents how widely the data set is spread along the direction which is perpendicular to the eigen vector.

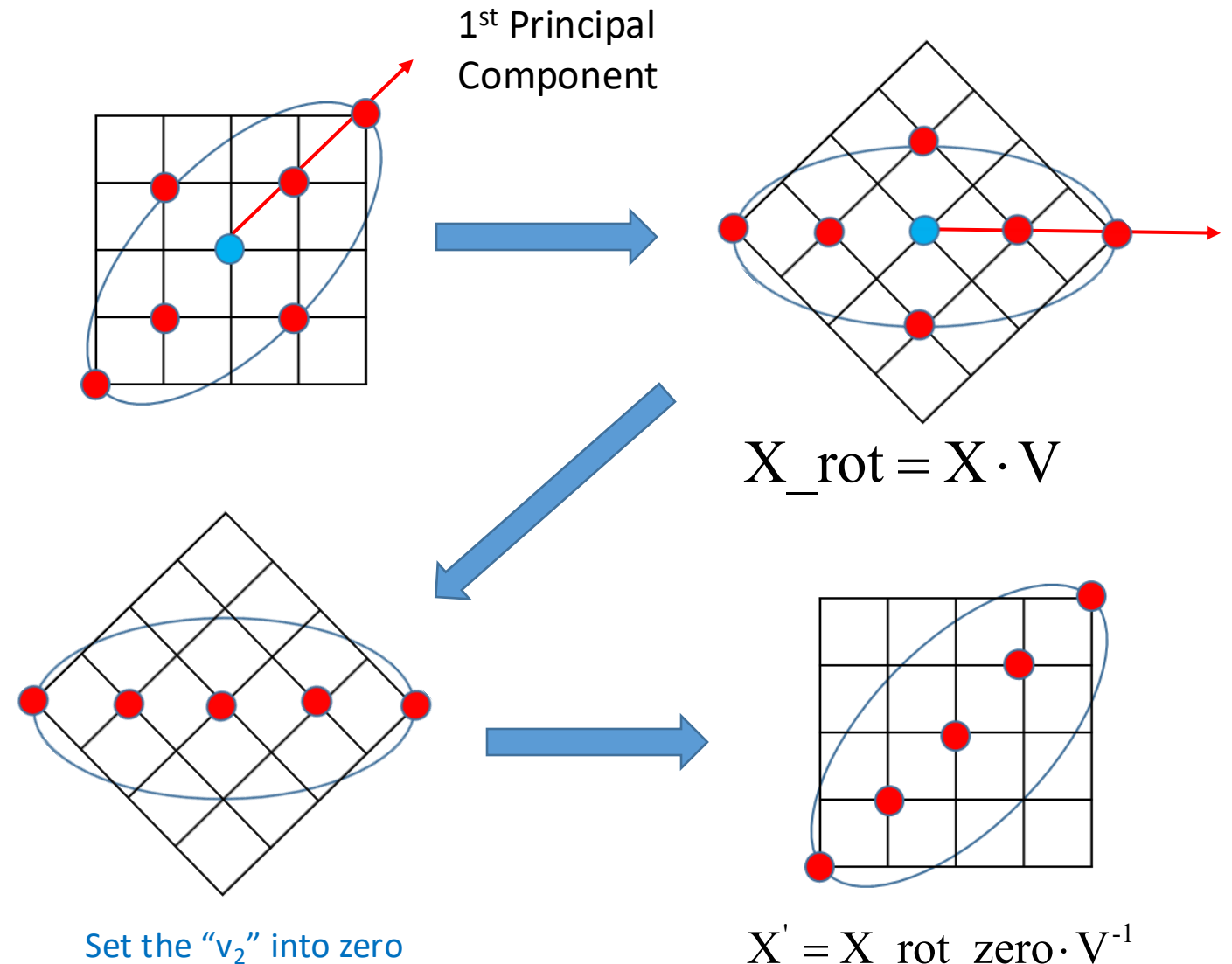
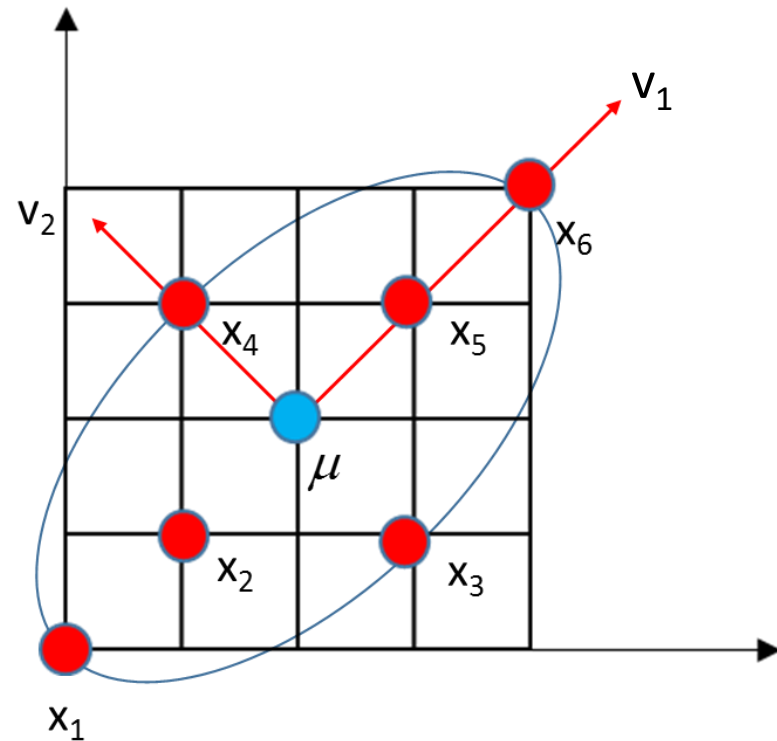


- ❑ A data point is defined by several, let's say, features,
- ❑ The number of features to define a data point is called the dimension of the data,
- ❑ High dimension data implies that it contains much information,
- ❑ Sometimes, we reduce its dimension, e.g., to visualize the data or to efficiently analyze them,
- ❑ PCA can reduce the dimension without losing relatively less information of the data.

- ❑ The previous example shows the case of two-dimensional data
- ❑ How can we reduce the two-dimensional data to one dimension?
- ❑ Yes, just project the data points onto the eigenvector space!



Dimension Reduction

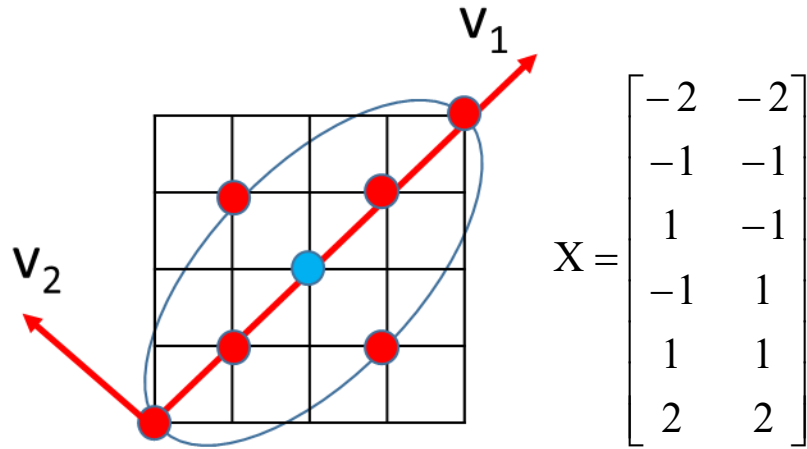


$$X_{\text{rot_zero}} = X_{\text{rot}} \cdot V^{-1}$$

$$X' = X_{\text{rot_zero}} \cdot V^{-1}$$

$$= X_{\text{rot_zero}} \cdot V^T$$

Dimension Reduction



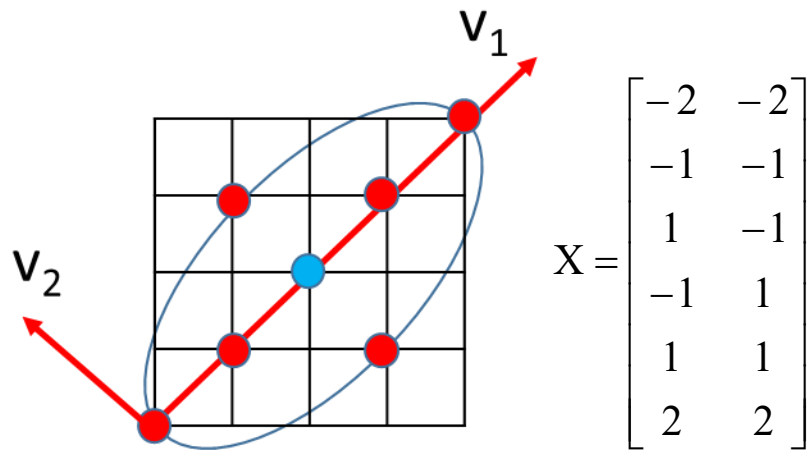
$$X = U\Sigma V^T$$

```
>> [U,S,V]=svd(x)
U =
-0.63246    0.00000    0.30819   -0.30819    0.28637    0.57274
-0.31623   -0.00000   -0.63635    0.63635    0.13426    0.26851
0.00000   -0.70711    0.50000    0.50000    0.00000    0.00000
-0.00000    0.70711    0.50000    0.50000   -0.00000   -0.00000
0.31623    0.00000   -0.00399    0.00399    0.94140   -0.11720
0.63246    0.00000   -0.00799    0.00799   -0.11720    0.76560

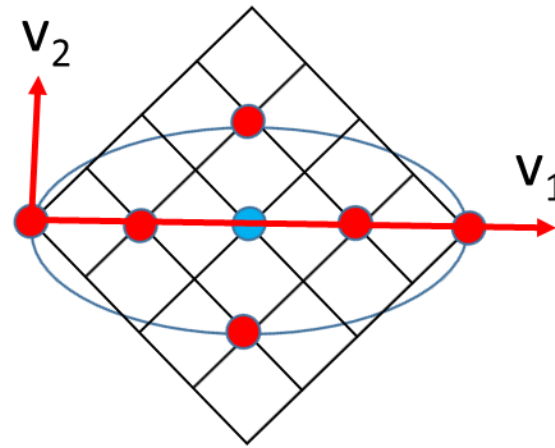
S =
Diagonal Matrix
4.4721     0
0    2.0000
0     0
0     0
0     0
0     0

V =
0.70711   -0.70711
0.70711    0.70711
```

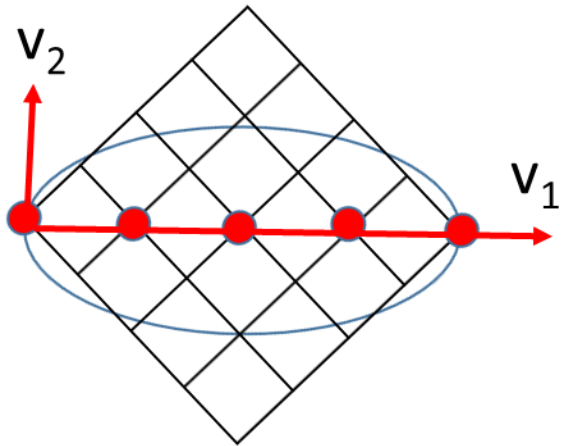
Dimension Reduction



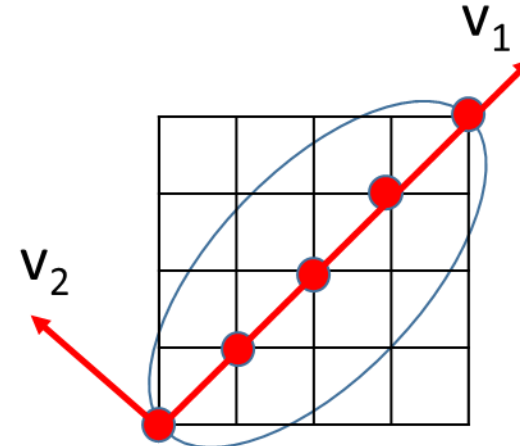
$$X = \begin{bmatrix} -2 & -2 \\ -1 & -1 \\ 1 & -1 \\ -1 & 1 \\ 1 & 1 \\ 2 & 2 \end{bmatrix}$$



$$X_{\text{rot}} = \begin{bmatrix} -2\sqrt{2} & 0 \\ -\sqrt{2} & 0 \\ 0 & -\sqrt{2} \\ 0 & \sqrt{2} \\ \sqrt{2} & 0 \\ 2\sqrt{2} & 0 \end{bmatrix} = \begin{bmatrix} -2 & -2 \\ -1 & -1 \\ 1 & -1 \\ -1 & 1 \\ 1 & 1 \\ 2 & 2 \end{bmatrix} \begin{bmatrix} \sqrt{2}/2 & -\sqrt{2}/2 \\ \sqrt{2}/2 & \sqrt{2}/2 \end{bmatrix}$$



$$X_{\text{rot_zero}} = \begin{bmatrix} -2\sqrt{2} & 0 \\ -\sqrt{2} & 0 \\ 0 & 0 \\ 0 & 0 \\ \sqrt{2} & 0 \\ 2\sqrt{2} & 0 \end{bmatrix}$$



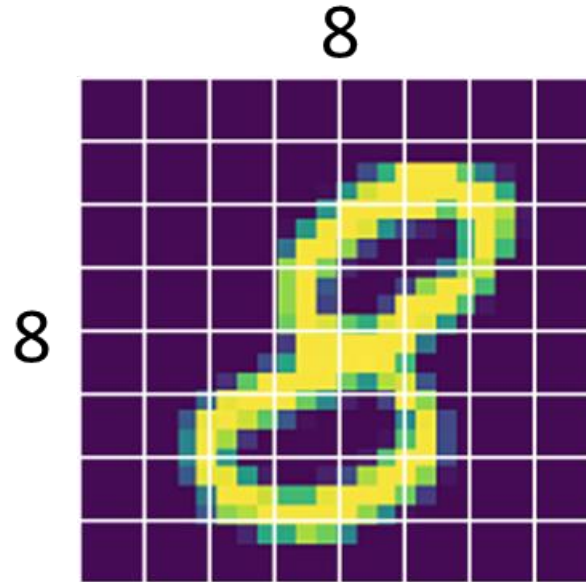
$$\begin{bmatrix} -2 & -2 \\ -1 & -1 \\ 0 & 0 \\ 0 & 0 \\ 1 & 1 \\ 2 & 2 \end{bmatrix}$$

Set the " v_2 " elements into zero

$$X' = X_{\text{rot_zero}} \cdot V^{-1} = X_{\text{rot_zero}} \cdot V^T$$

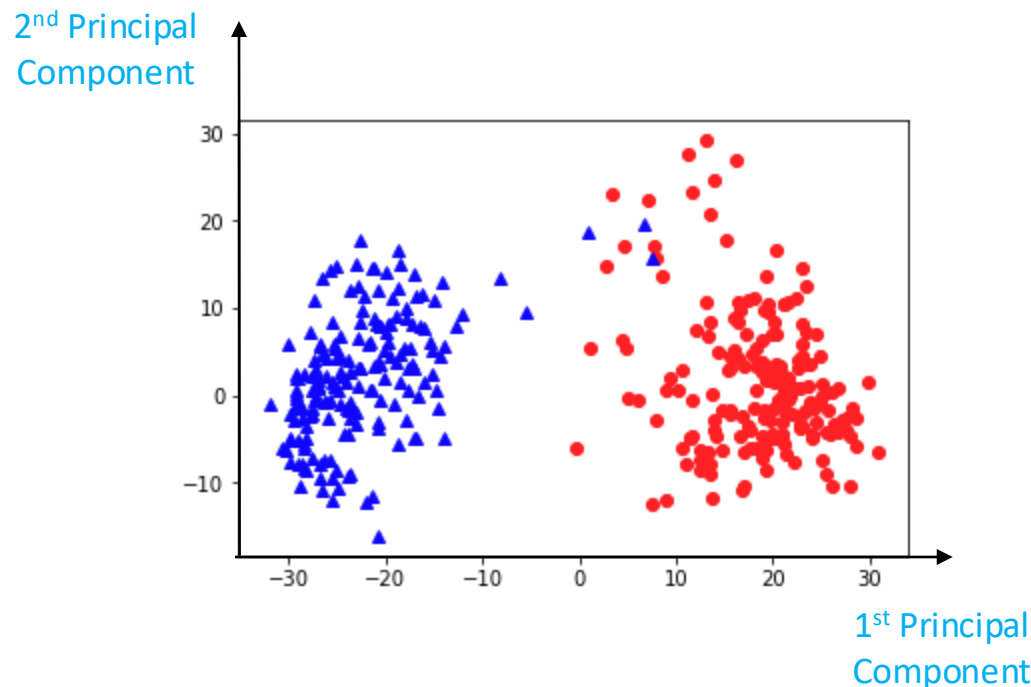
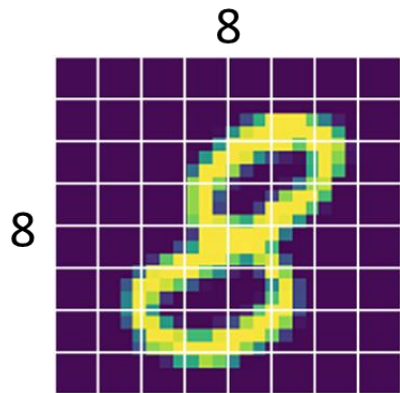
Dimension Reduction: an example

- ❑ Let's say, we have one image representing one data point as shown below,
- ❑ Then, we decide to present the data by all pixels which are 64 in this case, in other words, it is 64-dimensional data,
- ❑ What happens if we reduce its dimension to 2 dimension?



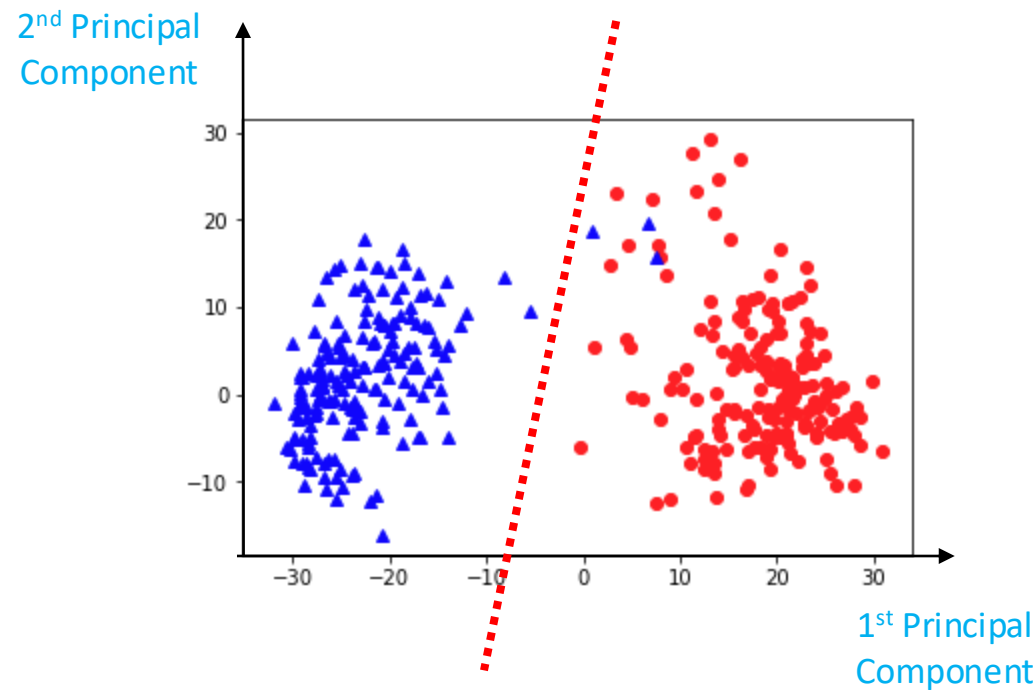
Dimension Reduction: an example

- ❑ Let's say, we have one image representing one data point as shown below,
- ❑ Then, we decide to present the data by all pixels which are 64 in this case, in other words, it is 64-dimensional data,
- ❑ What happens if we reduce its dimension to 2 dimension?



Dimension Reduction: an example

- ❑ Well, now we have a new set of data which have two dimension, so they can be presented in the two-dimensional space. Data visualization!
- ❑ Also, we may be able to classify those data by drawing a line???

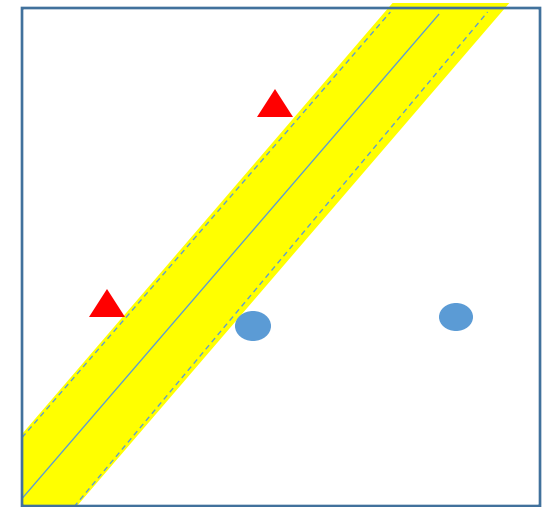
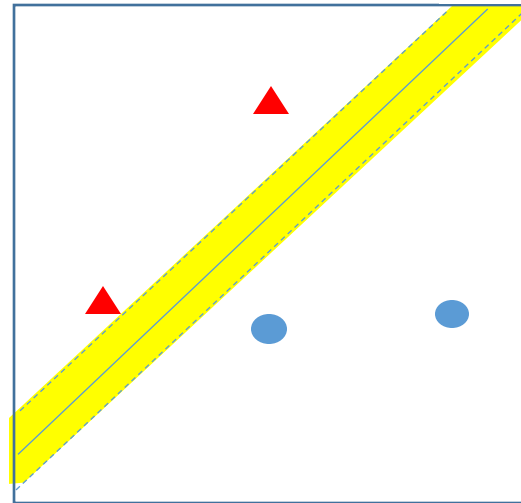
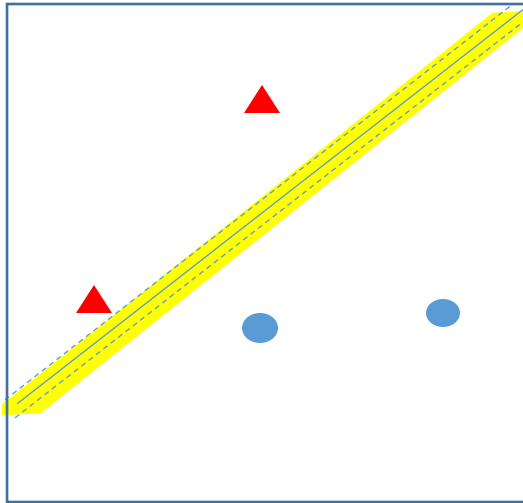


Support Vector Machine (SVM)

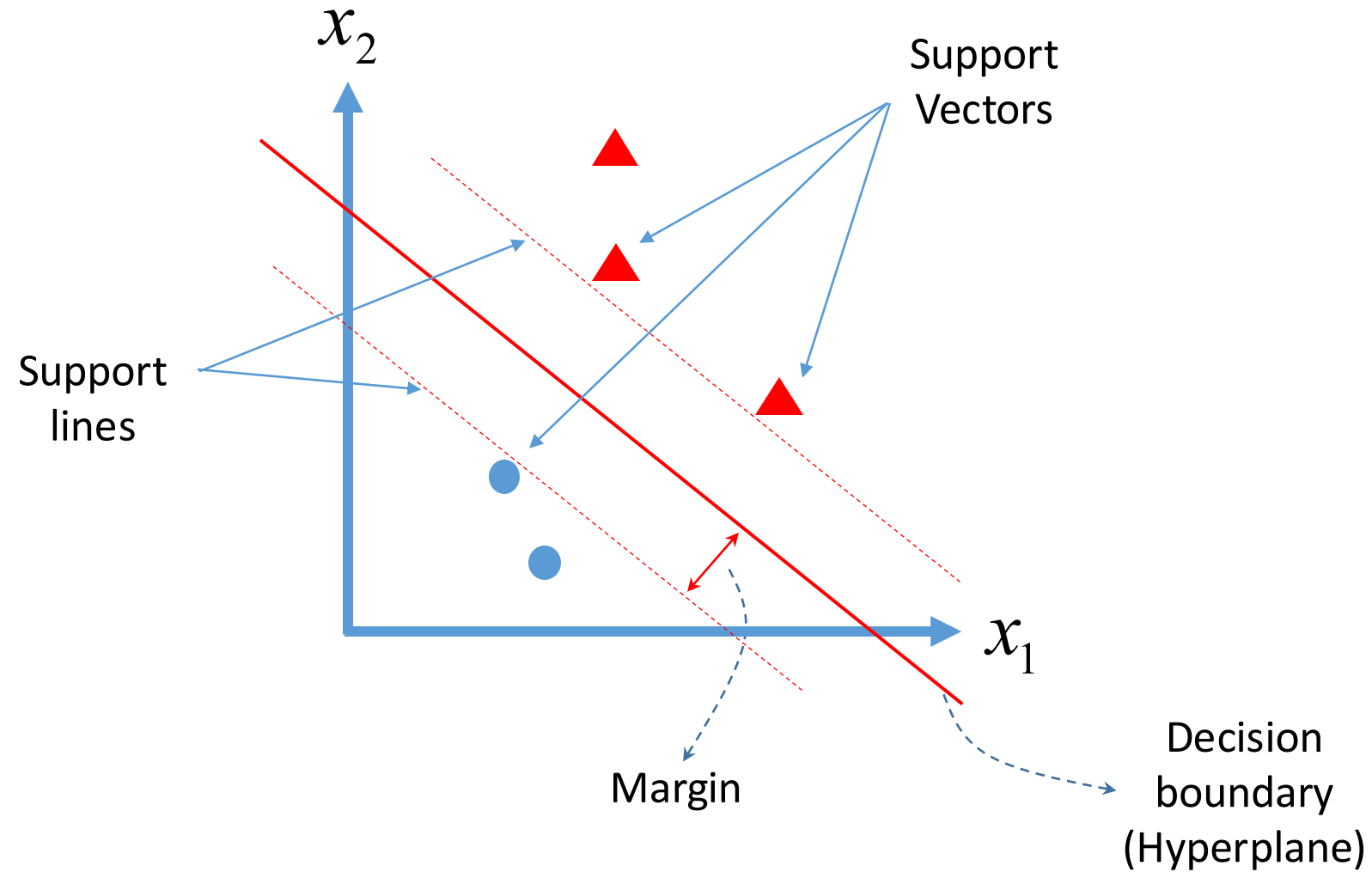
Why Support Vector Machine?

- ❑ Most widely used classification approach (practical)
 - Linearly separable data set
 - Non-linearly separable data set
- ❑ Supported by well defined mathematical theories
 - Geometry
 - Optimization

Which line is better to split two data sets?

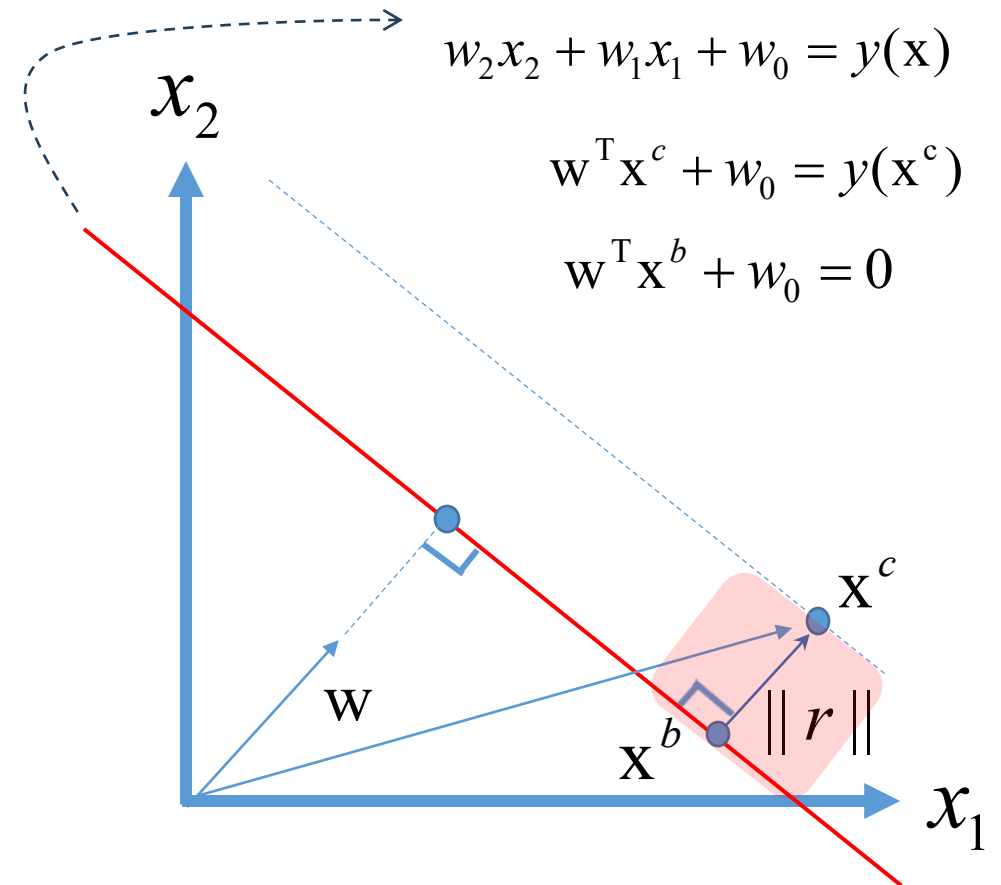


Terminology used in SVM



$$\mathbf{x}^c - \mathbf{x}^b = \underbrace{\|\mathbf{r}\|}_{\text{Margin}} \underbrace{\frac{\mathbf{w}}{\|\mathbf{w}\|}}_{\text{Unit vector showing the direction only}}$$

Size of the vector



Margin distance

$$\mathbf{x}^c = \mathbf{x}^b + \left\| \mathbf{r} \right\| \frac{\mathbf{w}}{\left\| \mathbf{w} \right\|}$$

□ Let's multiply $\mathbf{w}^T$ and add w_0 in both sides.

$$\mathbf{w}^T \mathbf{x}^c + w_0 = \mathbf{w}^T \mathbf{x}^b + w_0 + \mathbf{w}^T \left\| \mathbf{r} \right\| \frac{\mathbf{w}}{\left\| \mathbf{w} \right\|}$$

$$y(\mathbf{x}^c) = \mathbf{w}^T \left\| \mathbf{r} \right\| \frac{\mathbf{w}}{\left\| \mathbf{w} \right\|}$$

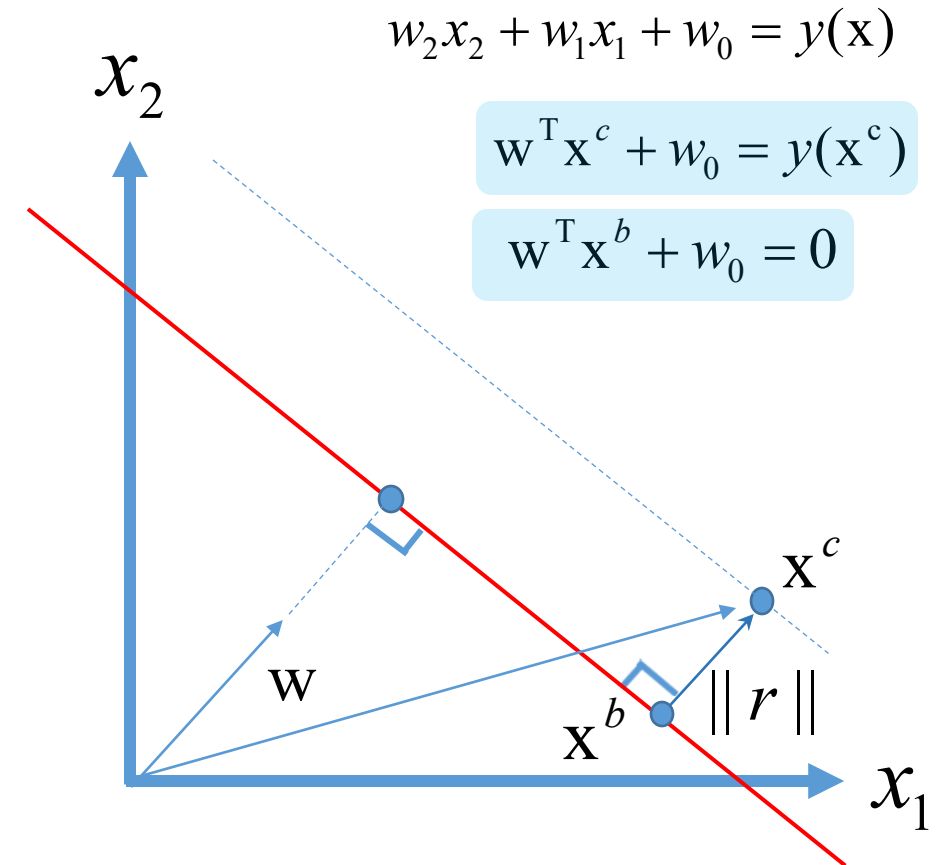
$$\left\| \mathbf{r} \right\| = \frac{y(\mathbf{x}^c)}{\left\| \mathbf{w} \right\|}$$

$$\left\| \mathbf{r} \right\| = \frac{1}{\left\| \mathbf{w} \right\|}$$

Let's say

$$|y(\mathbf{x}^c)| = 1$$

We use it later ...



Problem formulation

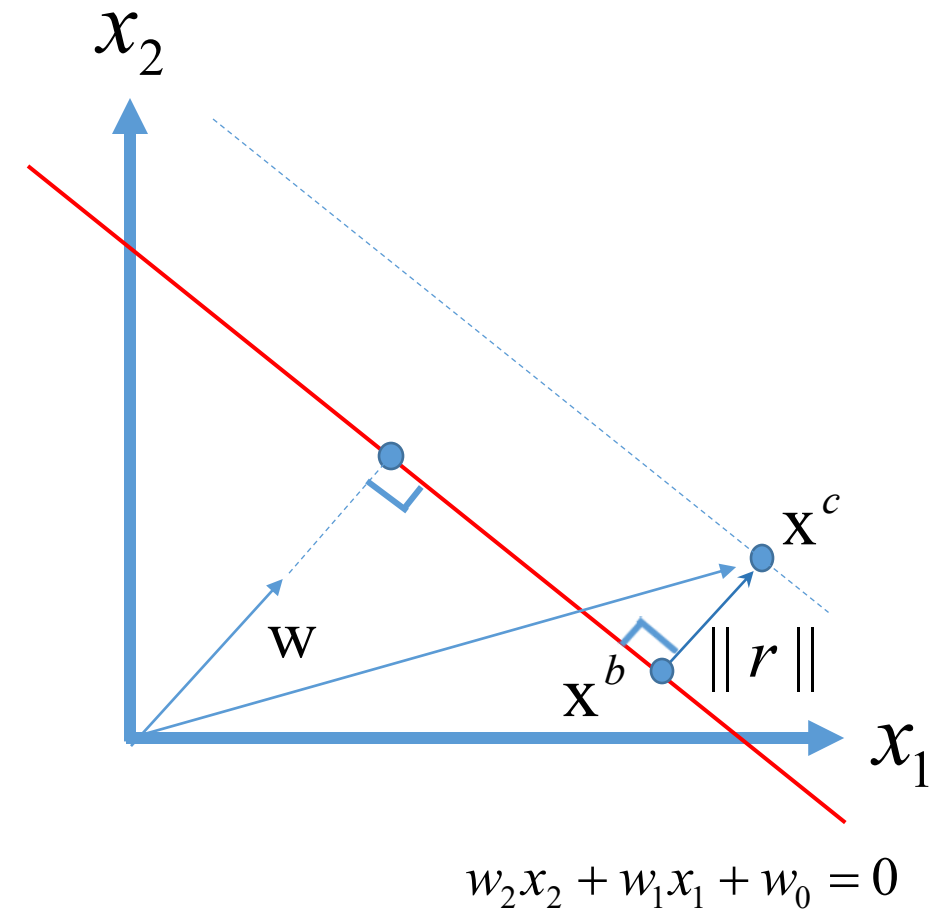
- Finding a decision boundary which maximizes the margin.

$$\max \|r\| = \frac{1}{\|w\|}$$

s.t.

$$t_n y(x_n) > 0 \quad \longrightarrow \quad \text{Every data points are classified correctly.}$$

$$\begin{cases} t_n = +1, & y(x_n) > 0 \\ t_n = -1, & y(x_n) < 0 \end{cases}$$



Problem formulation

- ❑ Let's modify the optimization problem a bit.

$$\max \frac{1}{\|w\|}$$

$$s.t. \quad t_n y(x_n) > 0, \quad \forall n$$

- ❑ Do you remember?

$$\max \frac{1}{\|w\|}$$

$$s.t. \quad t_n y(x_n) \geq 1, \quad \forall n$$

Let's say

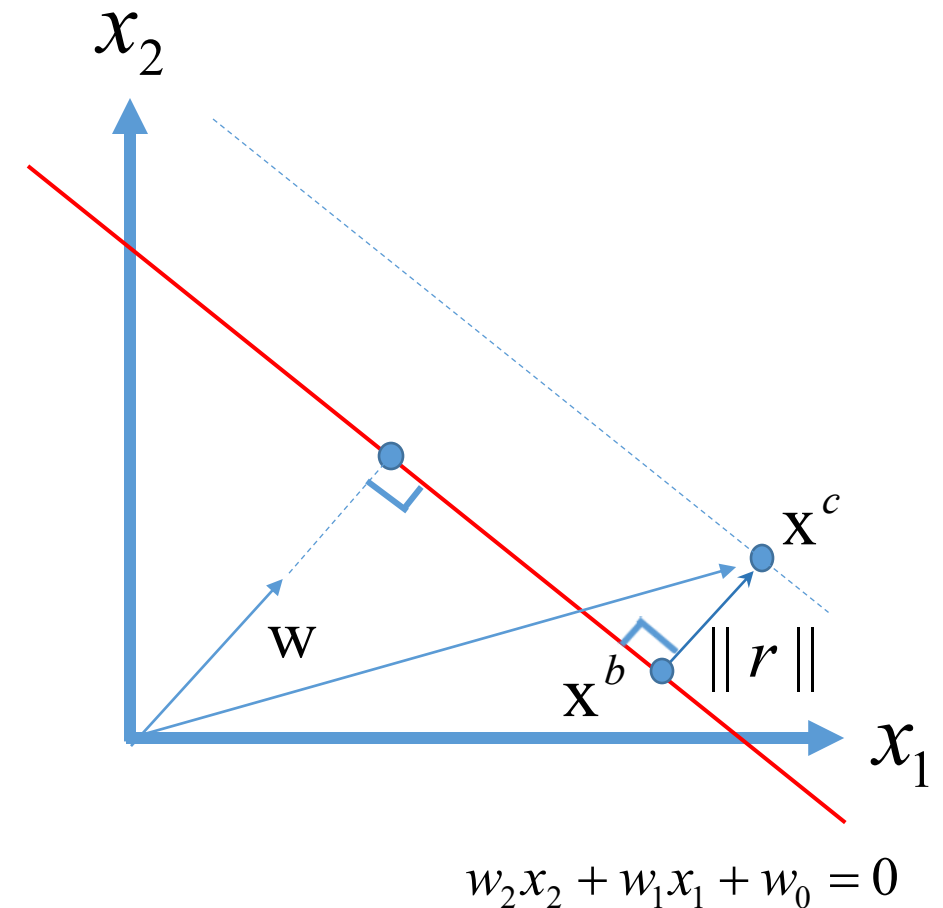
$$|y(x^c)| = 1$$

meaning that any data point is away from the decision boundary at least 1

- ❑ Finally

$$\min \frac{1}{2} \|w\|^2$$

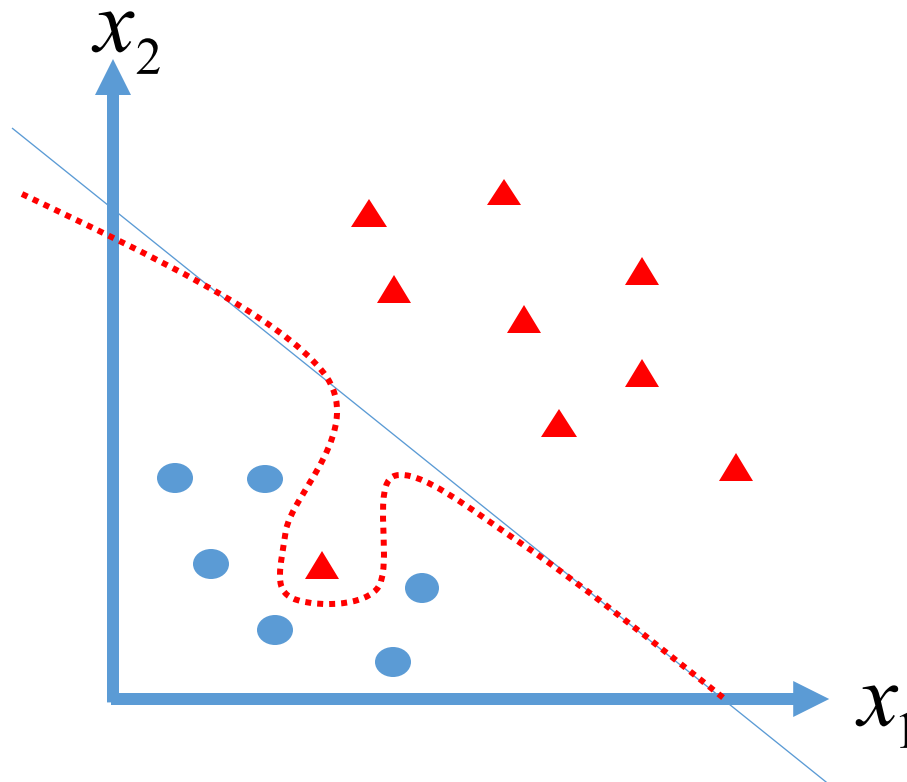
$$s.t. \quad t_n (w^T x_n + w_0) \geq 1, \quad \forall n$$



Quadratic programming

How about non-linearly separable case?

$$\begin{aligned} \min \quad & \frac{1}{2} ||w||^2 \\ \text{s.t.} \quad & t(x_n)(W \cdot x_n + w_0) \geq 1, \forall n \end{aligned}$$



Kernel Trick

Lagrange method for an optimization problem with inequality constraints

$$\min x^2$$

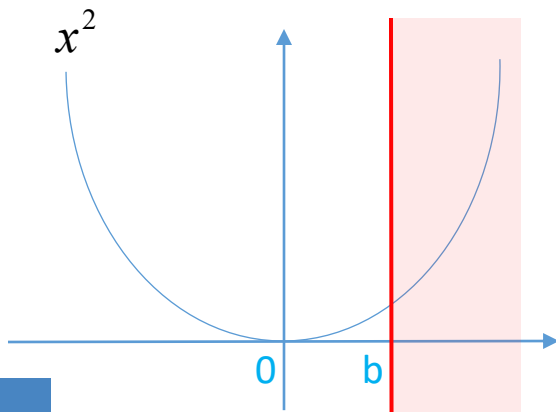
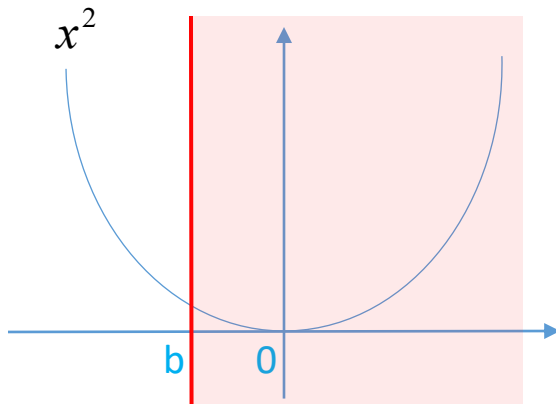
$$s.t. \quad x \geq b$$

?

=

$$\min_x \max_{\lambda} x^2 - \lambda(x - b)$$

$$s.t. \quad \lambda \geq 0$$



- ❑ If $b \leq 0$, the minima is 0 ... so $\lambda=0$
- ❑ If $b > 0$, the minima is b^2 ... so $x=b$
- ❑ So, either λ or $(x - b)$ becomes zero, in other words,
 - $\lambda(x-b) = 0$ (**complementary slackness**)
- ❑ Since $x \geq b$, maximizing λ minimizes the objective value
 - $\lambda \geq 0$

Convert the quadratic problem in SVM to Lagrange optimization problem

$$\min \frac{1}{2} \mathbf{w}^T \mathbf{w}$$

$$s.t. \quad t_n (\mathbf{w}^T \mathbf{x}_n + w_0) \geq 1$$



$$\min_{\mathbf{w}} \max_{\lambda} \frac{1}{2} \mathbf{w}^T \mathbf{w} - \sum_{n=1}^n \lambda_n (t_n (\mathbf{w}^T \mathbf{x}_n + w_0) - 1)$$

$$s.t. \quad \lambda_n \geq 0$$

Proof begins

Convert the quadratic problem in SVM to Lagrange optimization problem

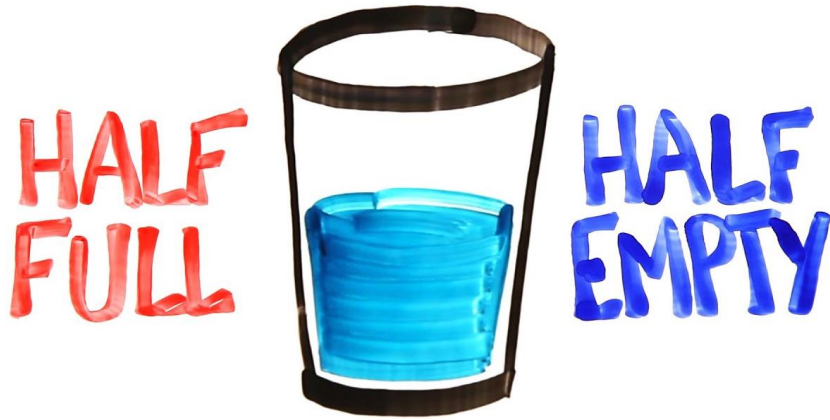
$$\begin{aligned} \min \quad & \frac{1}{2} \mathbf{w}^T \mathbf{w} \\ \text{s.t.} \quad & t_n (\mathbf{w}^T \mathbf{x}_n + w_0) \geq 1 \end{aligned}$$

$$\begin{aligned} \min_{\mathbf{w}} \max_{\lambda} \quad & \frac{1}{2} \mathbf{w}^T \mathbf{w} - \sum_{n=1}^n \lambda_n (t_n (\mathbf{w}^T \mathbf{x}_n + w_0) - 1) \\ \text{s.t.} \quad & \lambda_n \geq 0 \end{aligned}$$

- We would like to convert again the optimization problem above into another form, which provides same results.
- Because we want to solve the optimization problem in term of "lagrange multiplier (λ_n)".

$$\begin{aligned} \max_{\lambda} \min_{\mathbf{w}} \quad & \frac{1}{2} \mathbf{w}^T \mathbf{w} - \sum_{n=1}^n \lambda_n (t_n (\mathbf{w}^T \mathbf{x}_n + w_0) - 1) \\ \text{s.t.} \quad & \lambda_n \geq 0 \end{aligned}$$

Convert the quadratic problem in SVM to Lagrange optimization problem



Primal
problem



$$\begin{aligned} \min_w \max_{\lambda} \quad & \frac{1}{2} w^T w - \sum_{n=1}^n \lambda_n (t_n (w^T x_n + w_0) - 1) \\ \text{s.t.} \quad & \lambda_n \geq 0 \end{aligned}$$

- We would like to convert again the optimization problem above into another form, which provides same results.
 - Because we want to solve the optimization problem in term of "lagrange multiplier (λ_n)".

Dual
problem



$$\begin{aligned} \max_{\lambda} \min_w \quad & \frac{1}{2} w^T w - \sum_{n=1}^n \lambda_n (t_n (w^T x_n + w_0) - 1) \\ \text{s.t.} \quad & \lambda_n \geq 0 \end{aligned}$$

Convert the quadratic problem in SVM to Lagrange optimization problem

Karush–Kuhn–Tucker conditions

KKT conditions

1) Stationarity condition

$$\frac{\partial}{\partial \mathbf{w}} \frac{1}{2} \mathbf{w}^T \mathbf{w} - \frac{\partial}{\partial \mathbf{w}} \sum_{n=1}^n \lambda_n (t_n (\mathbf{w}^T \mathbf{x}_n + w_0) - 1) = 0$$

2) Complementary slackness condition

$$\lambda_n (t_n (\mathbf{w}^T \mathbf{x}_n + w_0) - 1) = 0$$

3) Duality feasibility condition

$$\lambda_n \geq 0$$

Primal
problem



$$\begin{aligned} \min_{\mathbf{w}} \max_{\lambda} & \frac{1}{2} \mathbf{w}^T \mathbf{w} - \sum_{n=1}^n \lambda_n (t_n (\mathbf{w}^T \mathbf{x}_n + w_0) - 1) \\ \text{s.t.} & \lambda_n \geq 0 \end{aligned}$$

- We would like to convert again the optimization problem above into another form, which provides same results.
 - Because we want to solve the optimization problem in term of "lagrange multiplier (λ_n)".

Dual
problem



$$\begin{aligned} \max_{\lambda} \min_{\mathbf{w}} & \frac{1}{2} \mathbf{w}^T \mathbf{w} - \sum_{n=1}^n \lambda_n (t_n (\mathbf{w}^T \mathbf{x}_n + w_0) - 1) \\ \text{s.t.} & \lambda_n \geq 0 \end{aligned}$$

Dual problem of the quadratic problem: applying stationarity condition

$$\max_{\lambda} \min_{w, w_0} L(w, w_0, \lambda) = \frac{1}{2} w^T w - \sum_{n=1}^N \lambda_n (t_n (w^T x_n + w_0) - 1)$$

$$w = \sum_{n=1}^N \lambda_n t_n x_n$$

$$\frac{\partial L}{\partial w} = w - \sum_{n=1}^N \lambda_n t_n x_n = 0$$

❖ The first one is called stationarity condition.

➤ when we partial differentiate the problem with respect to its parameter “w”, each of them should be zero.

Dual problem of the quadratic problem: applying stationarity condition

$$\max_{\lambda} \min_{w, w_0} L(w, w_0, \lambda) = \frac{1}{2} w^T w - \sum_{n=1}^N \lambda_n (t_n (w^T x_n + w_0) - 1)$$

$$w = \sum_{n=1}^N \lambda_n t_n x_n$$

$$\sum_{n=1}^N \lambda_n t_n = 0$$

$$\frac{\partial L}{\partial w} = w - \sum_{n=1}^N \lambda_n t_n x_n = 0$$

$$\frac{\partial L}{\partial w_0} = -\sum_{n=1}^N \lambda_n t_n = 0$$

❖ The first one is called stationarity condition.

➤ Again, this time in terms of “ w_0 ”

Dual problem of the quadratic problem: applying stationarity condition

$$\max_{\lambda} \min_{w, w_0} L(w, w_0, \lambda) = \frac{1}{2} w^T w - \sum_{n=1}^N \lambda_n (t_n (w^T x_n + w_0) - 1)$$

$$\boxed{w = \sum_{n=1}^N \lambda_n t_n x_n} \quad \begin{array}{c} \nearrow \\ \downarrow \end{array} \quad \boxed{\sum_{n=1}^N \lambda_n t_n = 0}$$

$$L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m x_n^T x_m - \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m x_n^T x_m - \sum_{n=1}^N \lambda_n t_n w_0 + \sum_{n=1}^N \lambda_n$$

Dual problem of the quadratic problem: applying stationarity condition

$$\max_{\lambda} \min_{w, w_0} L(w, w_0, \lambda) = \frac{1}{2} w^T w - \sum_{n=1}^N \lambda_n (t_n (w^T x_n + w_0) - 1)$$

$$w = \sum_{n=1}^N \lambda_n t_n x_n$$

$$\sum_{n=1}^N \lambda_n t_n = 0$$

$$L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m x_n^T x_m - \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m x_n^T x_m - \sum_{n=1}^N \lambda_n t_n w_0 + \sum_{n=1}^N \lambda_n$$

$$\max_{\lambda} L(\lambda) = \sum_{n=1}^N \lambda_n - \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m x_n^T x_m$$

Dual problem of the quadratic problem: applying stationarity condition

$$\max_{\lambda} \min_{w, w_0} L(w, w_0, \lambda) = \frac{1}{2} w^T w - \sum_{n=1}^N \lambda_n (t_n (w^T x_n + w_0) - 1)$$

$$w = \sum_{n=1}^N \lambda_n t_n x_n$$

$$\sum_{n=1}^N \lambda_n t_n = 0$$

$$L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m x_n^T x_m - \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m x_n^T x_m - \sum_{n=1}^N \lambda_n t_n w_0 + \sum_{n=1}^N \lambda_n$$

$$\max_{\lambda} L(\lambda) = \sum_{n=1}^N \lambda_n - \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m x_n^T x_m$$

$$\lambda_n \geq 0$$

$$\sum_{n=1}^N \lambda_n t_n = 0$$

$$w = \sum_{n=1}^N \lambda_n t_n x_n$$

"w" does not appear in the equation, and so we do not use this constraint anymore

Dual problem of the quadratic problem: applying stationarity condition

$$\begin{aligned} \max_{\lambda} L(\lambda) &= \sum_{n=1}^N \lambda_n - \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m \mathbf{x}_n^T \mathbf{x}_m \\ s.t. \quad \lambda_n &\geq 0, \quad \sum_{n=1}^N \lambda_n t_n = 0 \end{aligned}$$

- ❑ Let's change it to a quadratic programming again.
- ❑ As mentioned previously, a quadratic programming problem needs to be minimized

Dual problem of the quadratic problem: applying stationarity condition


$$\begin{aligned} \max_{\lambda} \quad & L(\lambda) = \sum_{n=1}^N \lambda_n - \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m \mathbf{x}_n^T \mathbf{x}_m \\ \text{s.t.} \quad & \lambda_n \geq 0, \quad \sum_{n=1}^N \lambda_n t_n = 0 \end{aligned}$$
$$\begin{aligned} \min_{\lambda} \quad & L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m \mathbf{x}_n^T \mathbf{x}_m - \sum_{n=1}^N \lambda_n \\ \text{s.t.} \quad & \lambda_n \geq 0, \quad \sum_{n=1}^N \lambda_n t_n = 0 \end{aligned}$$

□ Again, the optimization problem becomes a quadratic programming problem.

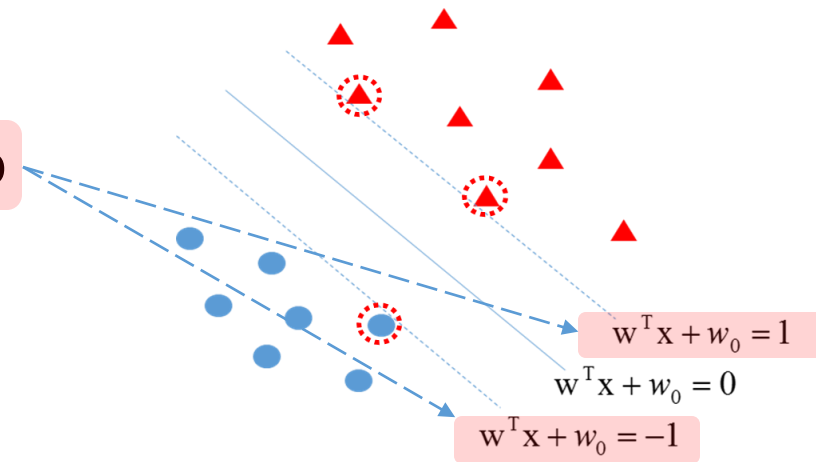
Let's summarize

$$\begin{aligned} \min_{\lambda} L(\lambda) &= \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m \mathbf{x}_n^T \mathbf{x}_m - \sum_{n=1}^N \lambda_n \\ \text{s.t. } \lambda &\geq 0, \quad t^T \lambda = 0 \end{aligned}$$

- ❑ The solution from the quadratic programming is “lagrange multipliers” (λ_n)
- ❑ Many of the solutions (lagrange multipliers) are zero
- ❑ Complementary slackness (one of KKT conditions) should be satisfied.

$$\lambda_n (t_n (w^T x_n + w_0) - 1) = 0$$

- ❑ In other words, if λ_n are not zero, $(t_n (w^T x_n + w_0) - 1)$ should be zero where corresponding data points should be support vectors.



Let's summarize

$$\min_{\lambda} L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m \mathbf{x}_n^T \mathbf{x}_m - \sum_{n=1}^N \lambda_n$$
$$s.t. \quad \lambda \geq 0, \quad t^T \lambda = 0$$

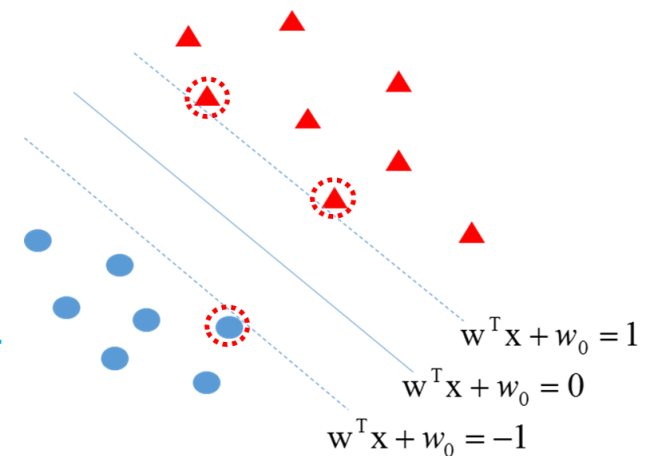
- ❑ The solution from the quadratic programming is “lagrange multipliers”(λ_n)
- ❑ Many of the solutions (lagrange multipliers) are zero
- ❑ Complementary slackness (one of KKT conditions) should be satisfied.

$$\lambda_n (t_n (w^T x_n + w_0) - 1) = 0$$

- ❑ In other words, if λ_n are not zero, $(t_n (w^T x_n + w_0) - 1)$ should be zero where corresponding data points should be support vectors.
- ❑ With the non-zero λ_n , w and w_0 can be calculated using $t_n (w^T x_n + w_0) = 1$

$$w = \sum_{n=1}^N \lambda_n t_n x_n$$

$$w_0 = t_n - \sum_{n=1}^N \lambda_n t_n x_n x_n$$



Let's summarize

$$\min_{\lambda} L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m \mathbf{x}_n^T \mathbf{x}_m - \sum_{n=1}^N \lambda_n$$

$$s.t. \quad \lambda \geq 0, \quad t^T \lambda = 0$$

We obtained previously

$$\begin{aligned} \square \quad & t_n(w_t x_n + w_0) = 1 \\ \Rightarrow \quad & (w_t x_n + w_0) = t_n \\ \Rightarrow \quad & w_0 = t_n - w_t x_n \end{aligned}$$

- ❑ The solution from the quadratic programming is "lagrange multipliers" (λ_n)
- ❑ Many of the solutions (lagrange multipliers) are zero
- ❑ Complementary slackness (one of KKT conditions) should be satisfied.

$$w_0 = t_n - \sum_{n=1}^N \lambda_n t_n x_n x_n$$

$$\lambda_n (t_n (w^T x_n + w_0) - 1) = 0$$

- ❑ In other words, if λ_n are not zero, $(t_n(w_t x_n + w_0) - 1)$ should be zero where corresponding data points should be support vectors.
- ❑ With the non-zero λ_n , w and w_0 can be calculated using $t_n(w_t x_n + w_0) = 1$

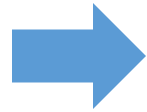
$$w = \sum_{n=1}^N \lambda_n t_n x_n$$

$$w_0 = t_n - \sum_{n=1}^N \lambda_n t_n x_n x_n$$

Proof ends

$$\min \frac{1}{2} \mathbf{w}^T \mathbf{w}$$

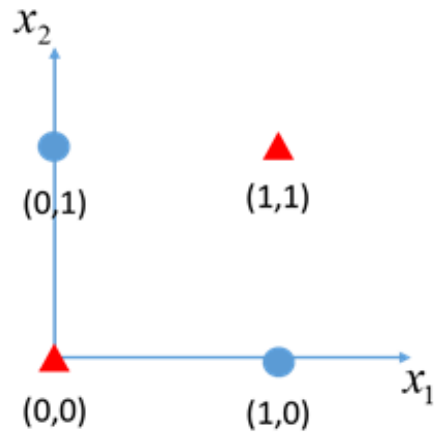
$$s.t. \quad t_n (\mathbf{w}^T \mathbf{x}_n + w_0) \geq 1$$



$$\min_{\lambda} L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m \mathbf{x}_n^T \mathbf{x}_m - \sum_{n=1}^N \lambda_n$$

$$s.t. \quad \lambda \geq 0, \quad t^T \lambda = 0$$

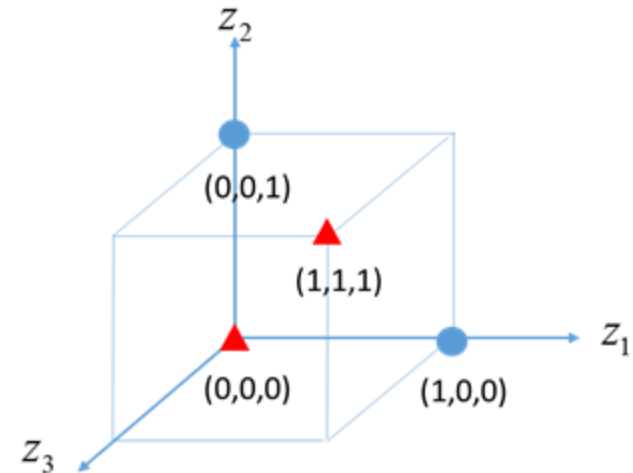
□ If data $\mathbf{x}_n$ are not linearly separable, what should we do?



Space X

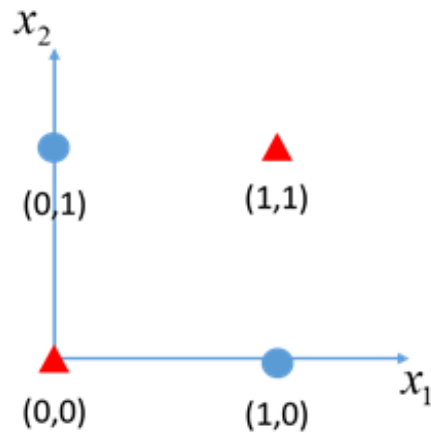
$$\phi(\mathbf{x}) = \begin{pmatrix} x_1^2 \\ x_1 x_2 \\ x_2^2 \end{pmatrix}$$

$$\phi(\mathbf{x}) \\ X \rightarrow Z$$

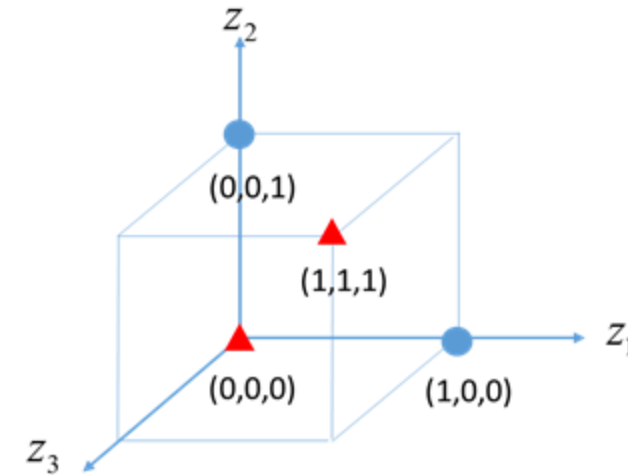
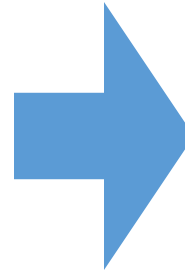


Space Z

- The idea of Kernel trick begins from here: to find the scalar values (the inner product of two vectors: $\mathbf{z}_n$ and $\mathbf{z}_m$) and so we can formulate the quadratic problem which can be linearly separable.



Space X



Space Z

$$\min_{\lambda} L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m \mathbf{x}_n^T \mathbf{x}_m - \sum_{n=1}^N \lambda_n$$

$$s.t. \quad \lambda \geq 0, \quad t^T \lambda = 0$$

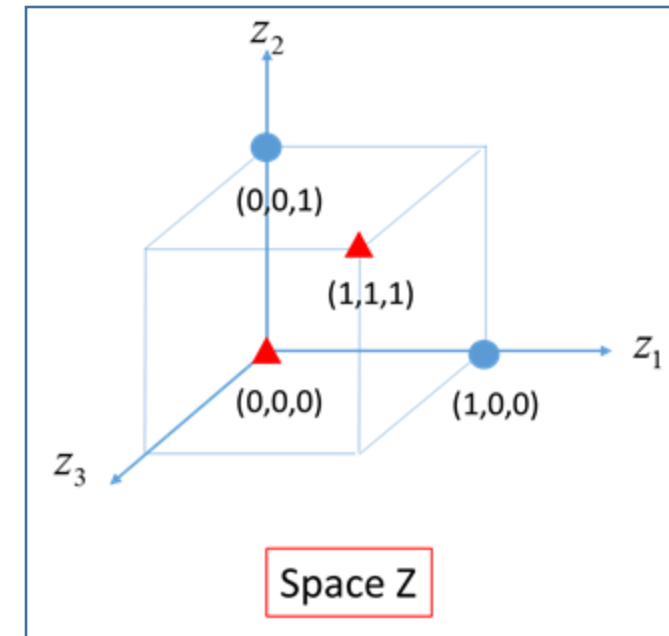
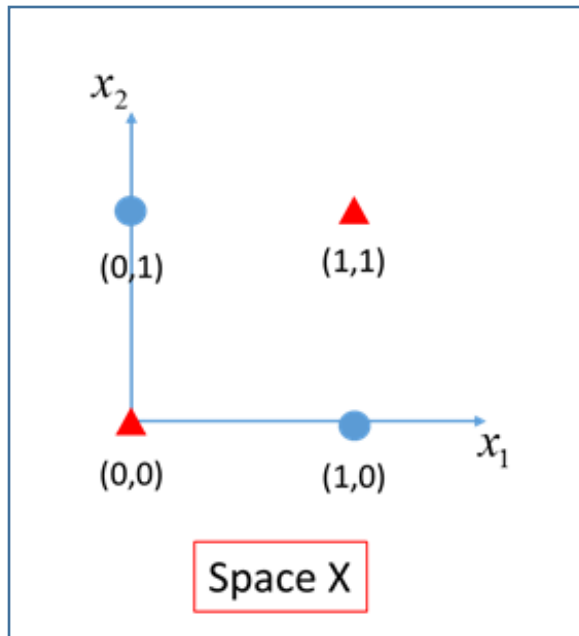
$$\min_{\lambda} L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m \mathbf{z}_n^T \mathbf{z}_m - \sum_{n=1}^N \lambda_n$$

$$s.t. \quad \lambda \geq 0, \quad t^T \lambda = 0$$

Kernel trick: Kernel function

- Kernel function $K()$ is a function which returns the scalar values (the inner product of two vectors: z_n and z_m in Z space) when the data points (x_n and x_m in X space) are given.

$$K(x_n^T, x_m) = z_n^T z_m$$



- With the Kernel function defined previously, we want to change the quadratic problem as follows:
- Because the Kernel function is a function of data points ($\mathbf{x}_n$ and $\mathbf{x}_m$) which we already have.

$$\min_{\lambda} L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m \mathbf{z}_n^T \mathbf{z}_m - \sum_{n=1}^N \lambda_n$$
$$s.t. \quad \lambda \geq 0, \quad t^T \lambda = 0$$



$$\min_{\lambda} L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m K(\mathbf{x}_n^T \mathbf{x}_m) - \sum_{n=1}^N \lambda_n$$
$$s.t. \quad \lambda \geq 0, \quad t^T \lambda = 0$$

Kernel trick: Kernel function

- With the Kernel function defined previously, we want to change the quadratic problem as follows:
- Because the Kernel function is a function of data points ($\mathbf{x}_n$ and $\mathbf{x}_m$) which we already have.

$$\min_{\lambda} L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m \mathbf{z}_n^T \mathbf{z}_m - \sum_{n=1}^N \lambda_n$$

$s.t. \quad \lambda \geq 0, \quad t^T \lambda = 0$

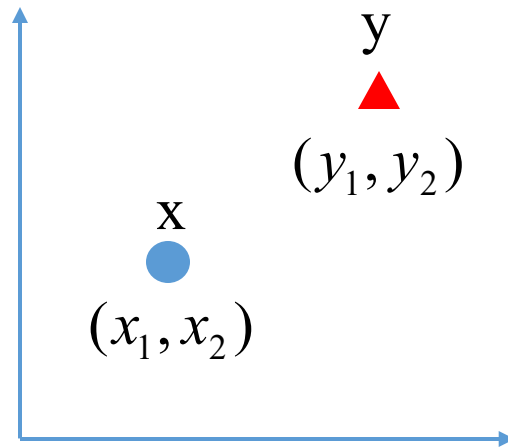


$$\min_{\lambda} L(\lambda) = \frac{1}{2} \sum_{n=1}^N \sum_{m=1}^N t_n t_m \lambda_n \lambda_m K(\mathbf{x}_n^T \mathbf{x}_m) - \sum_{n=1}^N \lambda_n$$

$s.t. \quad \lambda \geq 0, \quad t^T \lambda = 0$

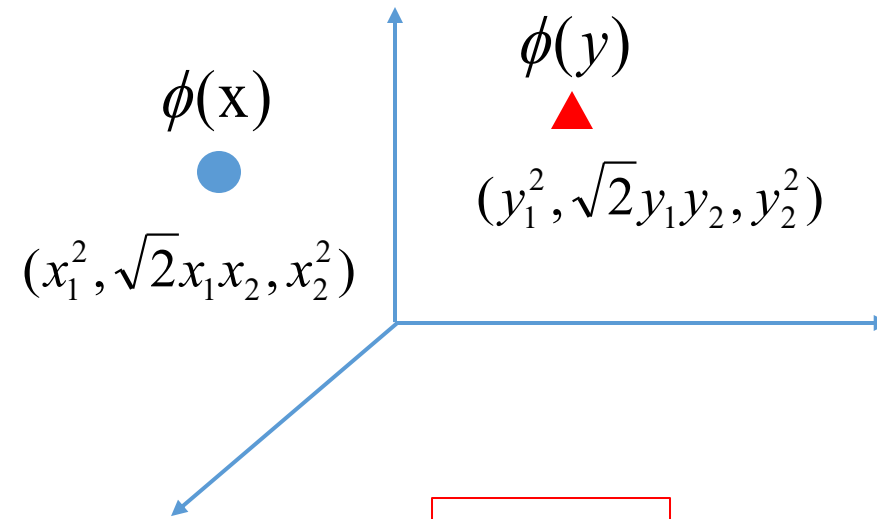
$$\min_{\lambda} L(\lambda) = \frac{1}{2} \lambda^T \begin{bmatrix} t_1 t_1 K(\mathbf{x}_1, \mathbf{x}_1) & t_1 t_2 K(\mathbf{x}_1^T, \mathbf{x}_2) & \cdots & t_1 t_N K(\mathbf{x}_1^T, \mathbf{x}_N) \\ t_2 t_1 K(\mathbf{x}_2, \mathbf{x}_1) & t_2 t_2 K(\mathbf{x}_2^T, \mathbf{x}_2) & \cdots & t_2 t_N K(\mathbf{x}_2^T, \mathbf{x}_N) \\ \cdots & \cdots & \cdots & \cdots \\ t_N t_1 K(\mathbf{x}_N, \mathbf{x}_1) & t_N t_2 K(\mathbf{x}_N^T, \mathbf{x}_2) & \cdots & t_N t_N K(\mathbf{x}_N^T, \mathbf{x}_N) \end{bmatrix} \lambda + (-1^T) \lambda$$

Polynomial kernel of degree 2



Space X

$$\begin{aligned} K(x, y) &= (xy)^2 \\ &= ((x_1, x_2) \cdot (y_1, y_2))^2 \\ &= (x_1 y_1 + x_2 y_2)^2 \\ &= x_1^2 y_1^2 + 2x_1 x_2 y_1 y_2 + x_2^2 y_2^2 \end{aligned}$$



Space Z

$$\begin{aligned} \phi(x)\phi(y) &= (x_1^2, \sqrt{2}x_1x_2, x_2^2) \cdot (y_1^2, \sqrt{2}y_1y_2, y_2^2) \\ &= x_1^2 y_1^2 + 2x_1 x_2 y_1 y_2 + x_2^2 y_2^2 \end{aligned}$$

- ❑ PCA and SVM are probably the most representative conventional machine learning algorithms.
- ❑ PCA helps you to manipulate a set of data in a way that
 - determining which features are important,
 - reducing its dimension, so that the data can be processed or visualized more efficiently.
- ❑ SVM is a classification method founded on well defined mathematical framework, which can handle linear or nonlinear classification problems.