

Chapter 6 Kernel method

We considered linear parametric models in which of the form of mapping $y(x, w)$ is governed by w of adaptive (trainable) parameters

Training set is used either to find a point estimate of w or to determine a posterior distribution of w

There is a class of ML in which the training data points are kept and used during the prediction phase

E.g. K-NN and Parzen probability density model

They require a metric to be defined that measures the similarity of any two vectors in input space

Linear parametric models can be written in a dual representation. In this form, predictions are made using linear combination of kernel functions, which are calculated using the training data points

This works well when the model uses a fixed non-linear mapping to transform the input data into a

feature space

$$K(x, x') = \Phi(x)^T \Phi(x')$$

The simplest kernel function comes from using the identity mapping for the feature transformation $\Phi(x) = x$

In this case,

$$K(x, x') = x^T x'$$

This called the linear kernel.

Kernel trick (kernel substitution)

When an algorithm uses input vectors only through dot products, replace the dot product with a kernel function to make it work in a nonlinear feature space

Dual representation

Consider a linear regression model with regularized SSE error function given by

$$J(w) := \frac{1}{2} \sum_{n=1}^N \{ w^T \Phi(x_n) - t_n \}^2 + \frac{\lambda}{2} w^T w \quad (6.2)$$

where $\lambda \geq 0$.

Set $\nabla_w J(w) = 0$. Then we see that the solution for

w takes the form of a linear combination of the

vectors $\Phi(x_n)$ with coefficients that are function of w

$$w = -\frac{1}{\lambda} \sum_{n=1}^N \{ w^T \Phi(x_n) - t_n \} \Phi(x_n) = \sum_{n=1}^N a_n \Phi(x_n) = \underbrace{\Phi^T}_{M \times N}^{N \times 1} a \quad (6.3)$$

where Φ is the design matrix and $\alpha = (\alpha_1, \dots, \alpha_N)^T$ with

$$\alpha_n = -\frac{1}{\lambda} \{ w_1^T \Phi(x_n) - t_n \}$$

We will reformulate the least squares algorithm in terms of the parameter vector α , which leads to a dual representation of the algorithm

Substitute $w_1 = \Phi^T \alpha$ into $J(w_1)$ (6.2)

$$J(\alpha) = \frac{1}{2} \alpha^T \Phi \Phi^T \Phi \Phi^T \alpha - \alpha^T \Phi \Phi^T \mathbf{t} + \frac{1}{2} \mathbf{t}^T \mathbf{t} + \frac{\lambda}{2} \alpha^T \Phi \Phi^T \alpha$$

where $\mathbf{t} = (t_1, t_2, \dots, t_N)^T$.

Now we define the gram matrix $K = \Phi \Phi^T$ which is an $N \times N$ symmetric matrix with elements

$$K_{nm} = \Phi(x_n)^T \Phi(x_m) = K(x_n, x_m)$$

In terms of the gram matrix, SSE error can be written as

$$J(\alpha) = \frac{1}{2} \alpha^T K \alpha - \alpha^T K \mathbf{z} + \frac{1}{2} \mathbf{z}^T \mathbf{z} + \frac{\lambda}{2} \alpha^T \alpha$$

Using (6.3) to eliminate w from (6.4) and solving for α we obtain

$$\alpha = (K + \lambda I_N)^{-1} \mathbf{z}$$

For linear regression model, we obtain the following prediction for a new input x

$$y(x) = w^T \Phi(x) = \alpha^T \Phi \Phi(x) = k(x)^T (\underbrace{K + \lambda I_N}_{\alpha})^{-1} \# \quad (6.9)$$

where $k(x)$ with $k_n(x) = K(x_n, x) = \Phi(x_n)^T \Phi(x)$

for $n = 1, 2, \dots, N$

The dual formulation allows the solution to the least square problem to be expressed by $K(x, x')$

Thus we see that the formulation allows the solution to the least square problem to be expressed entirely in

terms of the kernel function $k(x, x')$ (dual formulation)

α can be written as a linear combination of the feature vector $\Phi(x)$. So we can recover the original parameter vector w .

Remark

- To determine w , we need $O(M^3)$ calculations
- In the dual formulation, to determine α , we need $O(N^3)$
- The dual formulation uses only the kernel function, i.e. we do not need to compute the feature vector $\Phi(x)$ directly
- We can work in high (even infinite) dimensional space without extra cost (all through the kernel)

>1 $\frac{1}{2}$ basis function 방법

1. Fix a basis function $\Phi(x) = (\phi_0(x), \phi_1(x), \dots, \phi_{M-1}(x))^T$

2. Model $y(x) = w^T \Phi(x)$ with weight vector w

3. Finding w minimizing the cost function below

$$J(w) = \frac{1}{2} \sum_{n=1}^N \{ w^T \Phi(x_n) - t_n \}^2 + \frac{\lambda}{2} w^T w$$

where $\lambda \geq 0$.

kernel trick (dual representation)

1. Fix a kernel function $k(x, x')$
2. Model $y(x) = \mathbf{K}(x)^T \alpha$ where $\mathbf{K}(x) = (k_1(x), k_2(x), \dots, k_N(x))^T$
and $k_i(x) = k(x_i, x)$
3. Finding α minimizing the cost function below
$$J(\alpha) = \frac{1}{2} \alpha^T K K \alpha - \alpha^T K \mathbf{t} + \frac{1}{2} \mathbf{t}^T \mathbf{t} + \frac{\gamma}{2} \alpha^T K \alpha$$
where $K_{nm} = k(x_n, x_m)$, $\gamma \geq 0$

6.2 Constructing kernels

In order to use kernel substitution, we need to construct valid kernel functions.

One approach is to choose a feature space mapping $\Phi(x)$ and then use this to find the corresponding kernel

$$k(x, x') := \Phi(x)^T \Phi(x') = \sum_{i=1}^M \phi_i(x) \phi_i(x')$$

where $\phi_i(x)$ are basis functions.

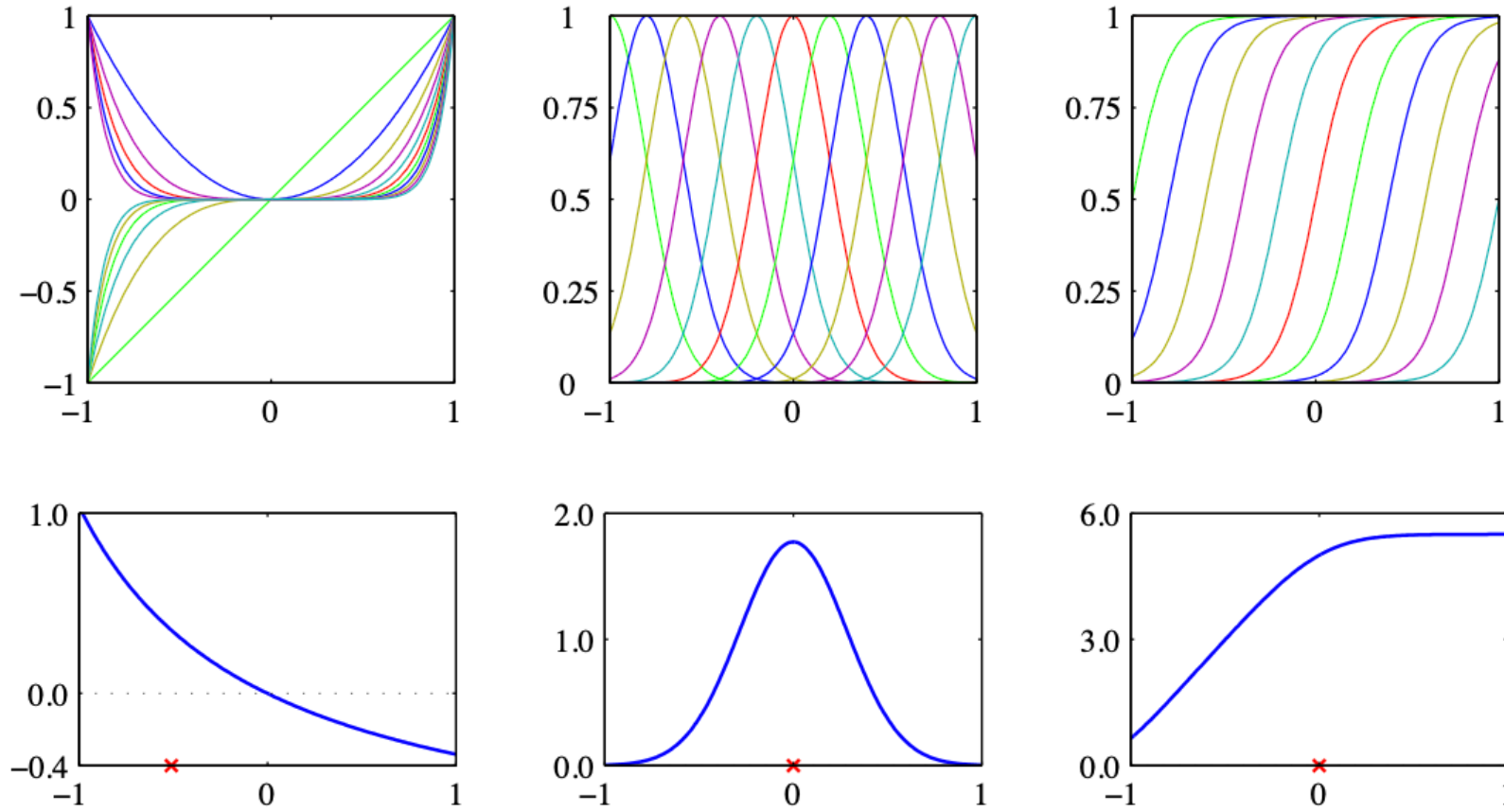


Figure 6.1 Illustration of the construction of kernel functions starting from a corresponding set of basis functions. In each column the lower plot shows the kernel function $k(x, x')$ defined by (6.10) plotted as a function of x , where x' is given by the red cross ($\times$), while the upper plot shows the corresponding basis functions given by polynomials (left column), 'Gaussians' (centre column), and logistic sigmoids (right column).

Alternative approach is to build kernel function directly without thinking about the feature mapping $\Phi(x)$

However, to be a valid kernel, the function must satisfy certain conditions: It should correspond to a dot product in some feature space (even infinite dimensional)

Mathematically, this means the kernel matrix K where

$K_{nm} = K(x_n, x_m)$ must be

- Symmetric $K_{nm} = K_{mn}$

- Positive semi-definite

E.g. Consider the kernel function given by

$$K(\mathbf{x}, \mathbf{z}) := (\mathbf{x}^T \mathbf{z})^2$$

In case of $D=2$, we can expand out the terms and identify the corresponding nonlinear feature mapping

$$\begin{aligned} K(\mathbf{x}, \mathbf{z}) &= (\mathbf{x}^T \mathbf{z})^2 = (x_1 z_1 + x_2 z_2)^2 \\ &= x_1^2 z_1^2 + 2x_1 z_1 x_2 z_2 + x_2^2 z_2^2 \end{aligned}$$

$$= (x_1^2, \sqrt{2} x_1 x_2, x_2^2) (\mathbf{z}_1^2, \sqrt{2} z_1 z_2, z_2^2)^T$$

$$= \Phi(\mathbf{x})^T \Phi(\mathbf{z})$$

Feature mapping $\Phi(\mathbf{x}) = (x_1^2, \sqrt{2} x_1 x_2, x_2^2)^T$

Here are the main properties that allow combining simpler kernels to build more powerful one:

Techniques for Constructing New Kernels.

Given valid kernels $k_1(\mathbf{x}, \mathbf{x}')$ and $k_2(\mathbf{x}, \mathbf{x}')$, the following new kernels will also be valid:

$$k(\mathbf{x}, \mathbf{x}') = ck_1(\mathbf{x}, \mathbf{x}') \quad (6.13)$$

$$k(\mathbf{x}, \mathbf{x}') = f(\mathbf{x})k_1(\mathbf{x}, \mathbf{x}')f(\mathbf{x}') \quad (6.14)$$

$$k(\mathbf{x}, \mathbf{x}') = q(k_1(\mathbf{x}, \mathbf{x}')) \quad (6.15)$$

$$k(\mathbf{x}, \mathbf{x}') = \exp(k_1(\mathbf{x}, \mathbf{x}')) \quad (6.16)$$

$$k(\mathbf{x}, \mathbf{x}') = k_1(\mathbf{x}, \mathbf{x}') + k_2(\mathbf{x}, \mathbf{x}') \quad (6.17)$$

$$k(\mathbf{x}, \mathbf{x}') = k_1(\mathbf{x}, \mathbf{x}')k_2(\mathbf{x}, \mathbf{x}') \quad (6.18)$$

$$k(\mathbf{x}, \mathbf{x}') = k_3(\phi(\mathbf{x}), \phi(\mathbf{x}')) \quad (6.19)$$

$$k(\mathbf{x}, \mathbf{x}') = \mathbf{x}^T \mathbf{A} \mathbf{x}' \quad (6.20)$$

$$k(\mathbf{x}, \mathbf{x}') = k_a(\mathbf{x}_a, \mathbf{x}'_a) + k_b(\mathbf{x}_b, \mathbf{x}'_b) \quad (6.21)$$

$$k(\mathbf{x}, \mathbf{x}') = k_a(\mathbf{x}_a, \mathbf{x}'_a)k_b(\mathbf{x}_b, \mathbf{x}'_b) \quad (6.22)$$

where $c > 0$ is a constant, $f(\cdot)$ is any function, $q(\cdot)$ is a polynomial with nonnegative coefficients, $\phi(\mathbf{x})$ is a function from $\mathbf{x}$ to $\mathbb{R}^M$, $k_3(\cdot, \cdot)$ is a valid kernel in $\mathbb{R}^M$, $\mathbf{A}$ is a symmetric positive semidefinite matrix, $\mathbf{x}_a$ and $\mathbf{x}_b$ are variables (not necessarily disjoint) with $\mathbf{x} = (\mathbf{x}_a, \mathbf{x}_b)$, and k_a and k_b are valid kernel functions over their respective spaces.

Polynomial kernel (example)

$$k(x, x') = (x^T x')^2$$

includes only terms of degree 2.

If we add a constant before squaring like

$$K(\mathbf{x}, \mathbf{x}') = (\mathbf{x}^T \mathbf{x}' + c)^2 \quad \text{where } c > 0$$

then the feature space becomes richer.

This kernel includes a constant (bias), linear terms (x_1, x_2) and degree 2 terms ($x_1 x_2$ or x_1^2)

$$K(\mathbf{x}, \mathbf{x}') = (x_1 x_1' + x_2 x_2' + c)^2 = (x_1 x_1')^2 + 2 x_1 x_1' x_2 x_2' + (x_2 x_2')^2$$

$$\mathbf{x} = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$$

$$+ 2c x_1 x_1' + 2c x_2 x_2' + c^2$$

$$= \Phi(\mathbf{x})^T \Phi(\mathbf{x}')$$

$$\Phi(\mathbf{x}) = (\sqrt{2c} x_1, \sqrt{2c} x_2, x_1^2, \sqrt{2} x_1 x_2, x_2^2, c)^T$$

In general,

$$k(x, x') = (x^T x')^M$$

gives a feature space with all monomials up to degree M

One commonly used kernel is the Gaussian kernel (a.k.a RBF kernel).

$$K(x, x') = \exp \left(-\frac{\|x - x'\|^2}{2\sigma^2} \right)$$

This kernel measures how close two vectors x and x' are.

$K(x, x')$ is close to 1 if x and x' are similar

$K(x, x')$ is close to 0 if x and x' are far apart

Although it looks like a Gaussian probability density, here it is not used as a probability

RBF kernel is a valid kernel

$$\|x - x'\|^2 = x^T x + (x')^T x' - 2x^T x'$$

So

$$K(x, x') = \exp\left(-\frac{x^T x}{2\sigma^2}\right) \exp\left(\frac{x^T x'}{\sigma^2}\right) \exp\left(-\frac{(x')^T x'}{2\sigma^2}\right)$$

Because of (6.14) and (6.16), together with the validity of the linear kernel $k(x, x') = x^T x'$

Remark

- The Gaussian kernel corresponds to an infinite-dimensional feature mapping $\Phi(x)$ (exercise 6.11)
- The Gaussian kernel is not restricted to the use of Euclidean distance. I.e. $x^T x'$ can be replaced with any nonlinear kernel $\tilde{K}(x, x')$,

$$K(x, x') = \exp \left\{ -\frac{1}{2\sigma^2} \left(\tilde{K}(x, x) + \tilde{K}(x', x') - 2\tilde{K}(x, x') \right) \right\}$$

One useful way to build kernel is by using a probabilistic generative model

Suppose we have a generative model $P(x)$. Then we can define a kernel as

$$K(x, x') = P(x) P(x')$$

This is a valid because it can be written as an inner product :

$$K(x, x') = \Phi(x)^T \Phi(x') \quad \text{where} \quad \Phi(x) = P(x)$$

Note that it is a very simple feature mapping, where each input x is just mapped to a scalar value $P(x)$.

In practice, more advanced versions use likelihood function or posterior distributions as features

It says that x and x' are similar if they both have high probabilities

Example: kernel from a Gaussian Mixture Model

Assume $p(x) = \sum_{k=1}^K \pi_k N(x | \mu_k, \Sigma_k)$

$$p(z=k) \quad p(x | z=k)$$

and estimates π_k, μ_k, Σ_k for $k=1, 2, \dots, K$

For each x , calculate the posterior responsibility

$$\phi_k(x) = p(z=k | x) = \frac{\pi_k N(x | \mu_k, \Sigma_k)}{\sum_{j=1}^K \pi_j N(x | \mu_j, \Sigma_j)}$$

Define $\Phi(x) := (\phi_1(x), \phi_2(x), \dots, \phi_K(x))^T$ and

$$k(x, x') = \Phi(x)^T \Phi(x')$$

This kernel measures the similarity between two vectors based on how responsible each Gaussian component is for generating them.

Kernels from mixture models with latent variable

We can build more flexible kernels by summing over multiple components in a probabilistic mixture model

1. Discrete case

Suppose we have

- latent variable $i \in \{1, 2, \dots, k\}$
- prob. distribution for each component $P(x | i)$
- prior over components $P(i)$

The kernel $k(x, x') = \sum_{i=1}^k P(x | i) P(x' | i) P(i)$

Remark

- This kernel is large when x and x' both have high likelihood under the same components
- It reflects how similarly the model explains the two inputs
- The latent variable i can be seen as a latent variable (e.g., cluster index in a mixture model)

2. Continuous case

If the latent variable z is continuous, we replace the sum with an integral

$$K(x, x') = \int p(x|z) p(x'|z) p(z) dz$$

where z is a continuous latent variable.

Another example of a kernel function is the sigmoidal kernel

$$k(\mathbf{x}, \mathbf{x}') = \tanh(a \mathbf{x}^T \mathbf{x}' + b)$$

Remark

- This function looks like the activation function in neural network

- This kernel is not always valid

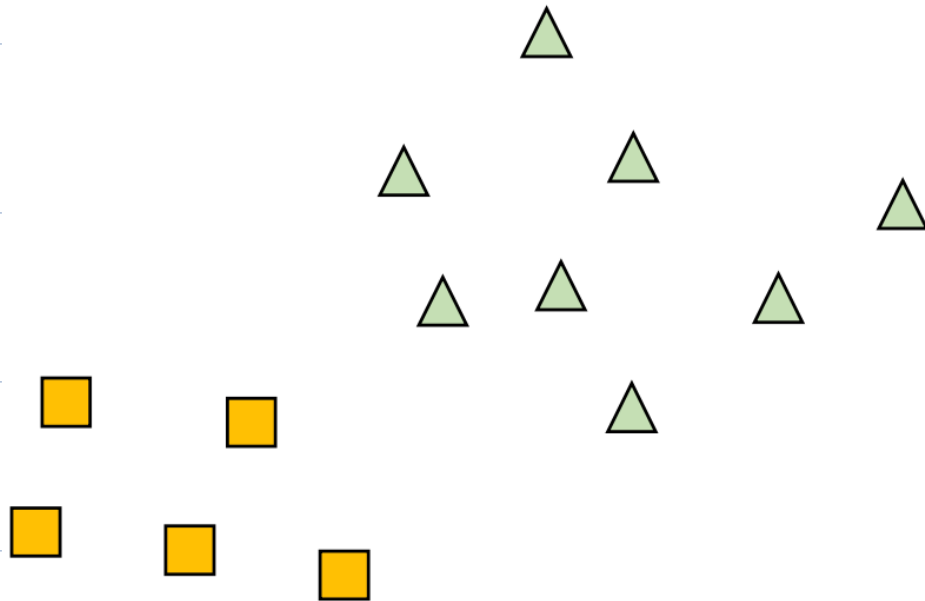
(its gram matrix is not guaranteed to be positive semidefinite for all a, b and data set).

Support vector machine

Two class classification problem

Linearly separable data set. $x \in \mathbb{R}^D$, $t = -1$ or 1

There are infinitely many decision boundaries. What is the best decision boundary?



SVM

$$w^T x + b$$

$$\min_{w} w^T w$$

quadratic
objective

$$\text{s.t.} \quad t_i (w^T x_i + b) \geq 1 \quad \forall i = 1, \dots, N$$

linear
constraints

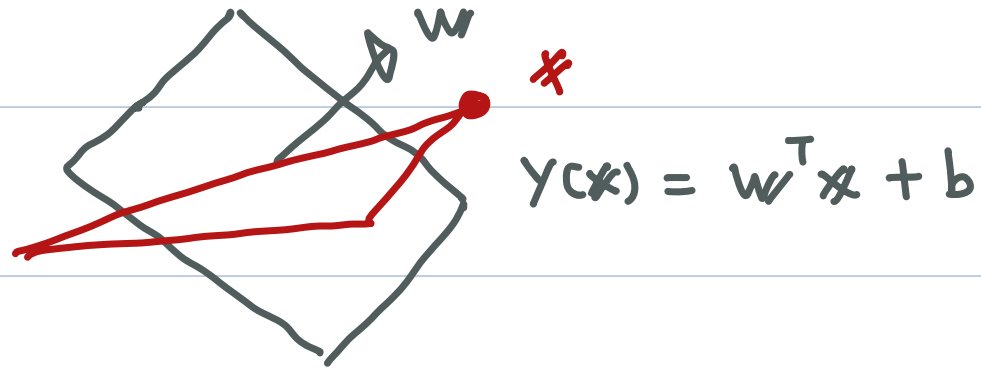
This is a quadratic optimization problem.

The one method that maximizes the distance to the
closest data points from both classes

- Maximal margin

Decision boundary

$y(x) = w^T x + b$ is a hyperplane



$$x = x_p + r \frac{w}{\|w\|}$$

$$y(x) = w^T x + b = w^T x_p + b + r \frac{w^T w}{\|w\|} = r \|w\|$$

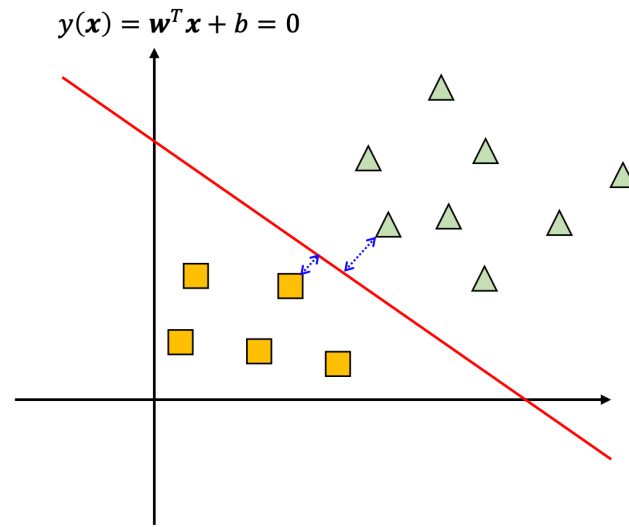
depends on x

Define a function $r: \mathbb{R}^D \rightarrow \mathbb{R}$ as

$$r(x) = \frac{y(x)}{\|w\|}$$

signed distance of x from the decision surface

Margin : The smallest distance between the decision boundary and any data points



w, b 으로
결정

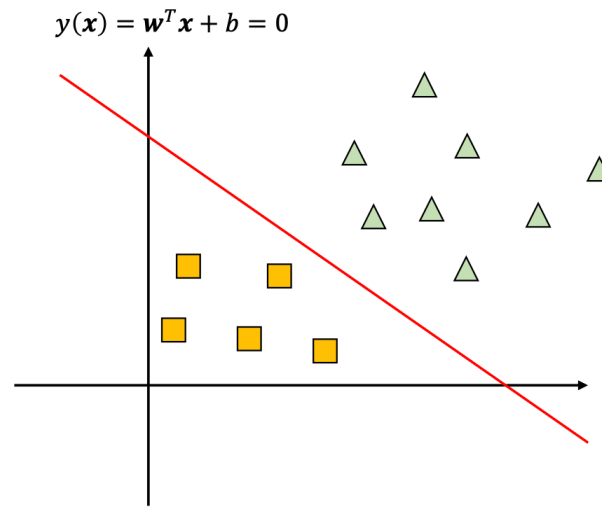
$$\delta(w, b) = \min_{x \in D} |r(x)| = \min_{x \in D} \frac{|y(x)|}{\|w\|} = \min_{x \in D} \frac{|w^T x + b|}{\|w\|}$$

Remark : $\delta(w, b)$ is scale invariant. I.e.,

$$\delta(\alpha w, \alpha b) = \delta(w, b) \quad \forall \alpha \neq 0$$

SVM : maximize margin

$$\max_{w, b} f(w, b) = \max_{w, b} \min_{x \in D} \frac{|w^T x + b|}{\|w\|}$$



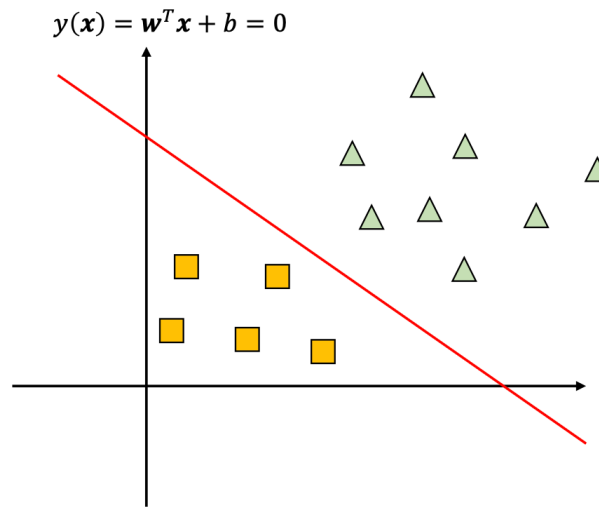
We need constraints

$$t_i (w^T x_i + b) \geq 0 \quad \forall i = 1, 2, \dots, N$$

I.e., max margin classifier with constraints

$$\max_{w, b} f(w, b) = \max_{w, b} \min_{x \in D} \frac{|w^T x + b|}{\|w\|}$$

$$\text{s.t.} \quad t_i (w^T x_i + b) \geq 0 \quad \forall i = 1, 2, \dots, N$$



Quite complicated. Need to simplify the problem

1. $\|w\|$ does not depend on x

$$\Rightarrow \max_{w, b} \min_{x \in D} \frac{|w^T x + b|}{\|w\|} = \max_{w, b} \frac{1}{\|w\|} \left(\min_{x \in D} |w^T x + b| \right)$$

2. $\delta(w, b)$ is scale invariant

$\Rightarrow$ so we can add a constraint

$$\min_{x \in D} |w^T x + b| = 1$$

3. 'max' can be replaced by 'min'

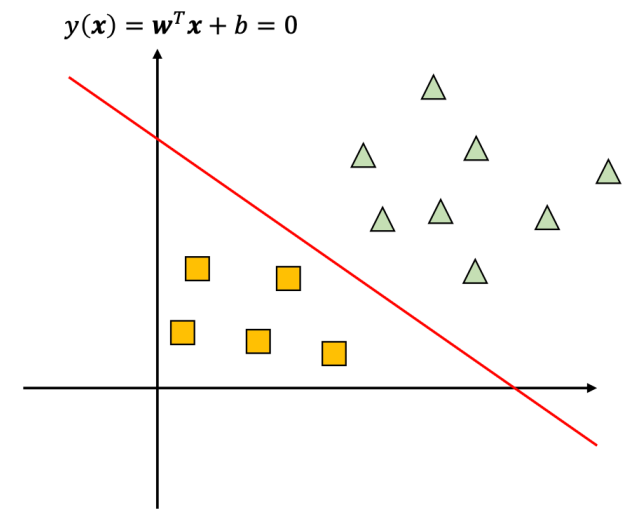
$$\max_{w, b} \frac{1}{\|w\|} \Rightarrow \min_{w, b} \|w\|^2$$

Max margin classifier

$$\min_{w, b} \frac{1}{2} \|w\|^2$$

$$\text{s.t.} \quad t_i (w^T x_i + b) \geq 1 \quad \forall i=1, \dots, N$$

$$\min_i |w^T x_i + b| = 1$$



We need to simplify the constraints

Max margin classifier

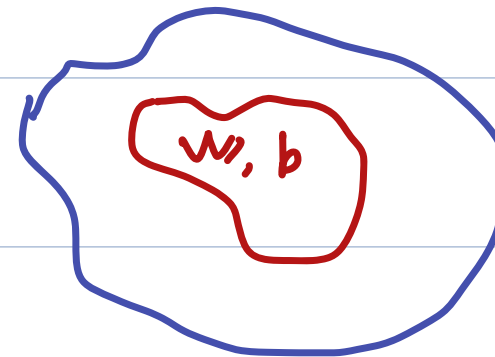
$$\min_{w, b} \frac{1}{2} \|w\|^2 \quad \text{s.t.} \quad t_i (w^T x_i + b) \geq 0 \quad \forall i=1, \dots, N$$

$$\min_i |w^T x_i + b| = 1 \quad \textcircled{1}$$

$$\Leftrightarrow \min_{w, b} \frac{1}{2} \|w\|^2 \quad \text{s.t.} \quad t_i (w^T x_i + b) \geq 1 \quad \forall i=1, \dots, N \quad \textcircled{2}$$

constraints 를 만족하는

w, b 집합



Assume w_* , b_* is the optimal solution and

$$\min_i t_i (w_*^T x_i + b_*) = \alpha > 1$$

Let $\tilde{w} = \frac{1}{\alpha} w_*$ and $\tilde{b} = \frac{1}{\alpha} b_*$. Then $(\tilde{w}, \tilde{b})$ is

a feasible solution i.e. $t_i (\tilde{w}^T x_i + \tilde{b}) \geq 1 \quad \forall i = 1, \dots, N.$

and $\|\tilde{w}\| = \frac{1}{\alpha} \|w_*\| < \|w_*\|$

It contradicts with the assumption that is w_* optimal.

$\Rightarrow$ Optimal solution must have $t_{\hat{i}} (w^T x_{\hat{i}} + b) = 1$ for some $\hat{i}$

SVM

$$\min_{w, b} \frac{1}{2} w^T w$$

quadratic
objective

$$\text{s.t.} \quad t_i (w^T x_i + b) \geq 1 \quad \forall i = 1, \dots, N$$

linear
constraints

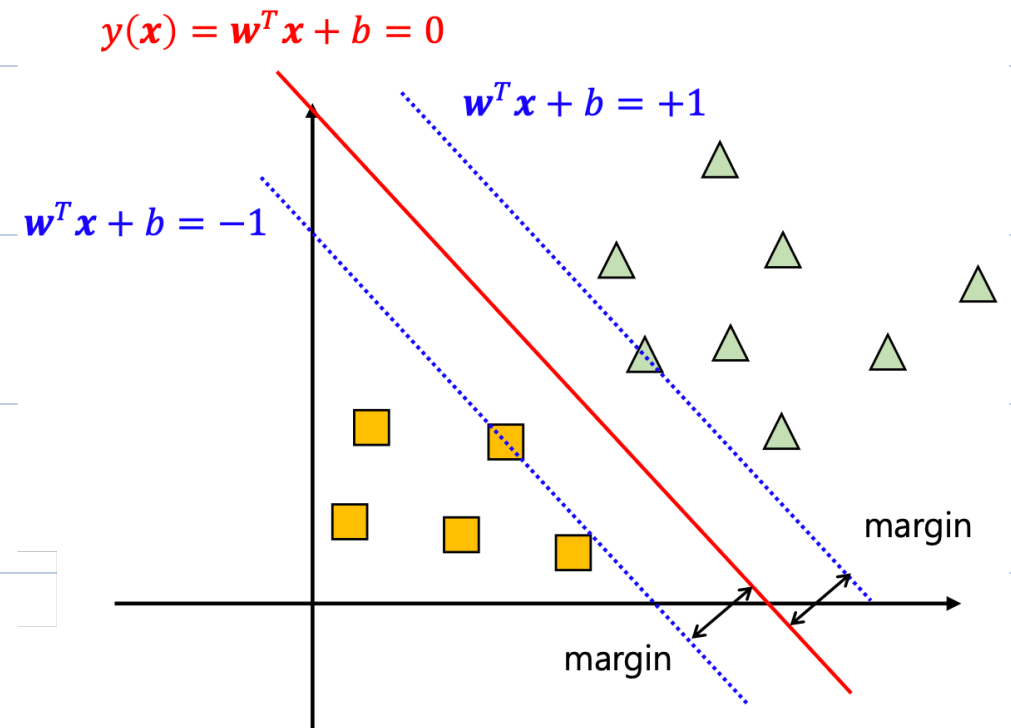
Convex problem

Support vector

$$- t_i (w^T x_i + b) = 1.$$

Decision boundary

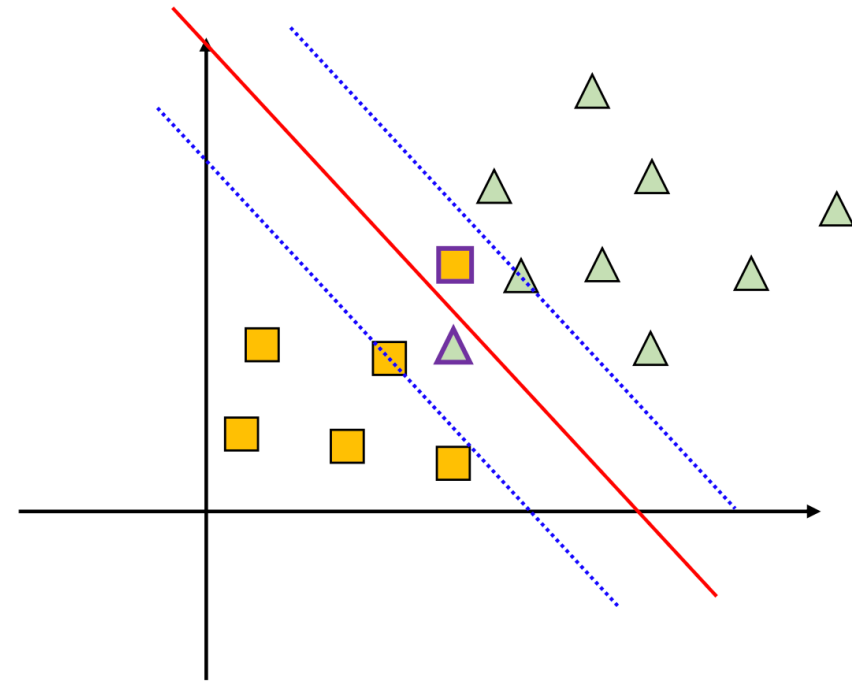
- lies in the middle of the region separating data points



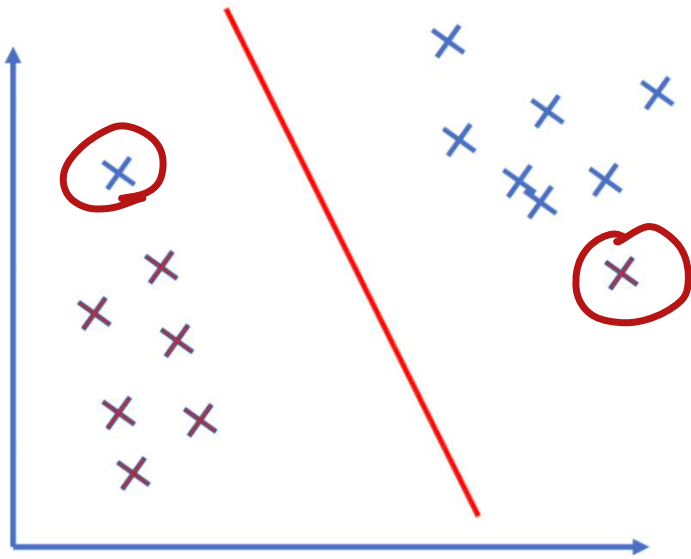
Support vector machine

Overlapping classes

- Non-linearly separable
- Soft constraints
- Soft margin



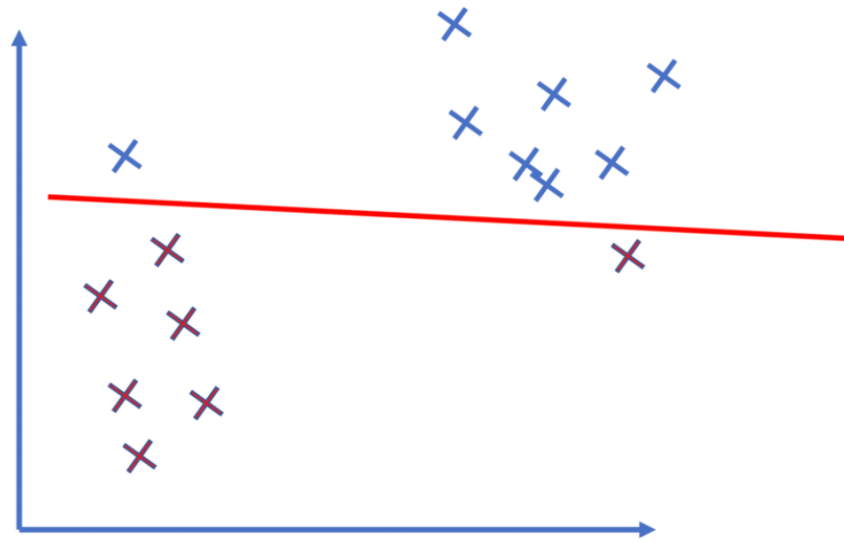
Which one is better?



It looks OK.

But train error
is not zero.

$$t_{\hat{\lambda}}(w^T x_{\hat{\lambda}} + b) \neq 1 \quad \forall i$$



Train error is zero.

But over-fitted

SVM

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

$$\text{s.t.} \quad t_i (w^T x_i + b) \geq 1 \quad \forall i = 1, \dots, N$$

$\Downarrow$

$$\min_{w, \xi} \frac{1}{2} \|w\|^2 + C \cdot \sum_{i=1}^N \xi_i$$

$$\text{s.t.} \quad t_i (w^T x_i + b) \geq 1 - \xi_i$$

$$\forall i = 1, \dots, N$$

$$\xi_i \geq 0$$

Slack variable ξ

- Data points are allowed to be on 'wrong side' of the margin boundary but with a penalty that increases with the distance from that boundary

$$\min_{w, \xi} \quad \frac{1}{2} \|w\|^2 + C \cdot \sum_{i=1}^N \xi_i$$

$$\text{s.t.} \quad t_i (w^T x_i + b) \geq 1 - \xi_i$$

$$\xi_i \geq 0$$

$$\Phi(x)^T \Phi(x') \quad \text{or} \quad x^T x'$$



$$\underline{K(x, x')}$$

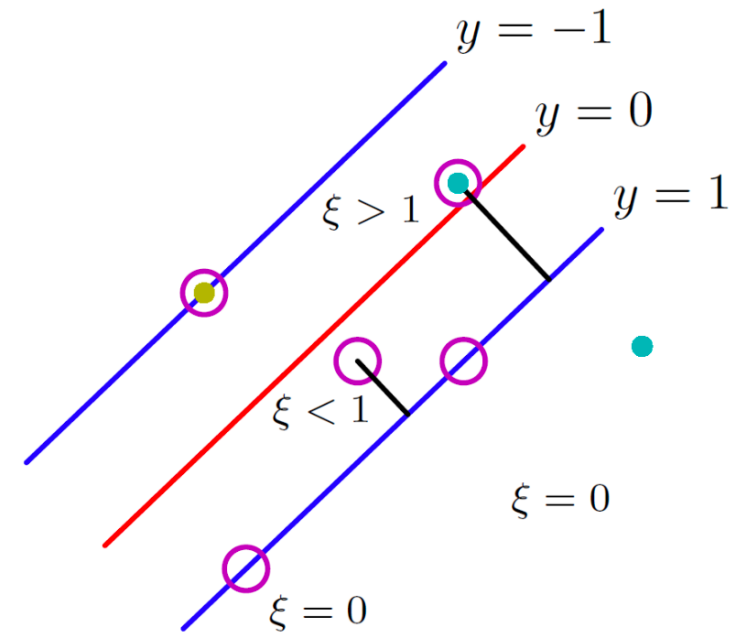
$$\forall i = 1, \dots, N$$

$\xi_i = 0$: correctly classified

$0 < \xi_i \leq 1$: inside the margin

but on the correct side

$\xi_i > 1$: on the wrong side of the
decision boundary



$$\min_{w, \xi} \frac{1}{2} \|w\|^2 + C \cdot \sum_{i=1}^N \xi_i$$

$$\text{s.t.} \quad t_i (w^T x_i + b) \geq 1 - \xi_i$$

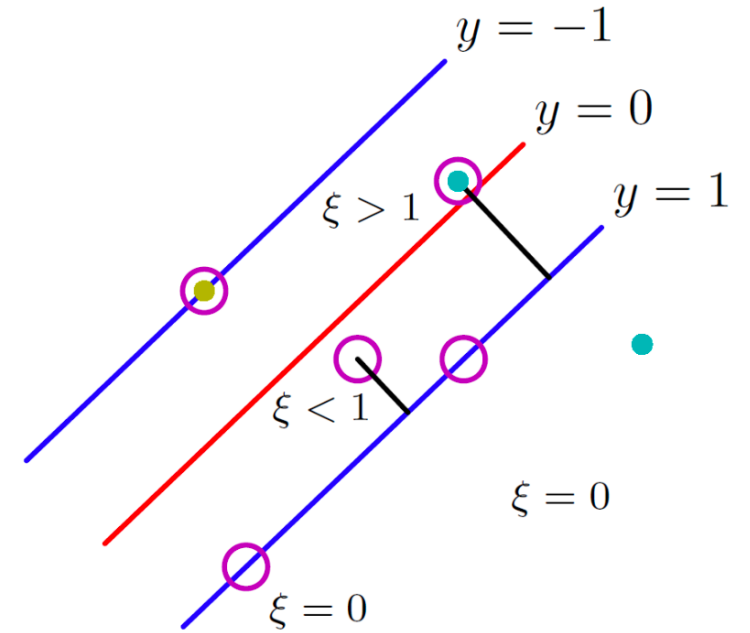
$$\xi_i \geq 0$$

$$i = 1, \dots, N$$

If C is large, then $\xi \rightarrow 0$.

SVM tries to classify all training data correctly

Overfitting and weak regularization



$$\min_{w, \xi} \frac{1}{2} \|w\|^2 + C \cdot \sum_{i=1}^N \xi_i$$

$$\text{s.t.} \quad t_i (w^T x_i + b) \geq 1 - \xi_i$$

$$\xi_i \geq 0$$

$$\forall i = 1, \dots, N$$

Overlapping classes : Hinge loss

$$\min_{w, b} \frac{1}{2} \|w\|^2 + c \cdot \sum_{i=1}^N \xi_i$$

$\phi^T \phi' \rightarrow K(\phi, \phi')$

$$\text{s.t.} \quad t_i (w^T \phi_i + b) \geq 1 - \xi_i$$

$$\forall i=1, \dots, N$$

$$\xi_i \geq 0$$

$$\Rightarrow \xi_i \geq 0 \quad \text{and} \quad \xi_i \geq 1 - t_i \overbrace{(w^T \phi_i + b)}^{\gamma(\phi_i)}$$

$$\Rightarrow \xi_i := \max(0, 1 - t_i \gamma(\phi_i))$$

$\Rightarrow$

$$\min_{w, \xi} \quad \frac{1}{2} \|w\|^2 + c \cdot \sum_{\bar{i}=1}^N \xi_{\bar{i}}$$

$$\text{s.t.} \quad \xi_{\bar{i}} = \max(0, 1 - t_{\bar{i}} \gamma(x_{\bar{i}})) \quad \forall \bar{i}=1, \dots, N$$

$$\text{where} \quad \gamma(x_{\bar{i}}) = w^T x_{\bar{i}} + b$$

$\Rightarrow$

$$\min_{w, \xi} \quad \frac{1}{2} \|w\|^2 + c \cdot \sum_{\bar{i}=1}^N \max(0, 1 - t_{\bar{i}} \gamma(x_{\bar{i}}))$$

or

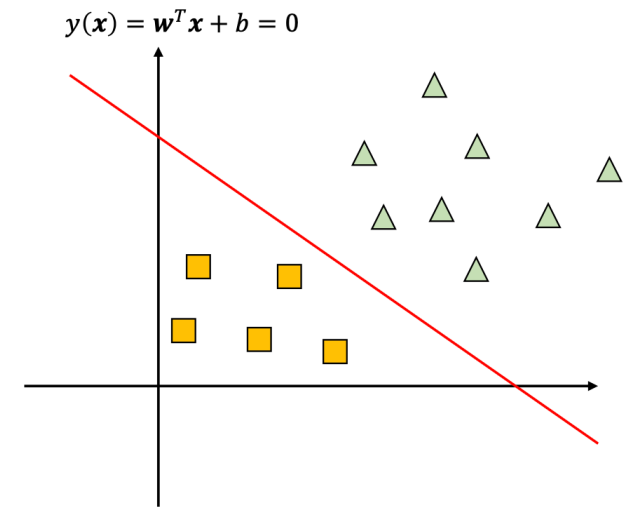
$$\min_{w, \xi} \quad \sum_{\bar{i}=1}^N \max(0, 1 - t_{\bar{i}} \gamma(x_{\bar{i}})) + \frac{1}{2c} \|w\|^2$$

$$\text{where} \quad \gamma(x_{\bar{i}}) = w^T x_{\bar{i}} + b$$

SVM = Linear classifier with Hinge loss

$$\min_{w, b} \sum_{i=1}^N \max(0, 1 - t_i y(x_i)) + \frac{1}{2c} \|w\|^2$$

where $y(x_i) = w^T x_i + b$



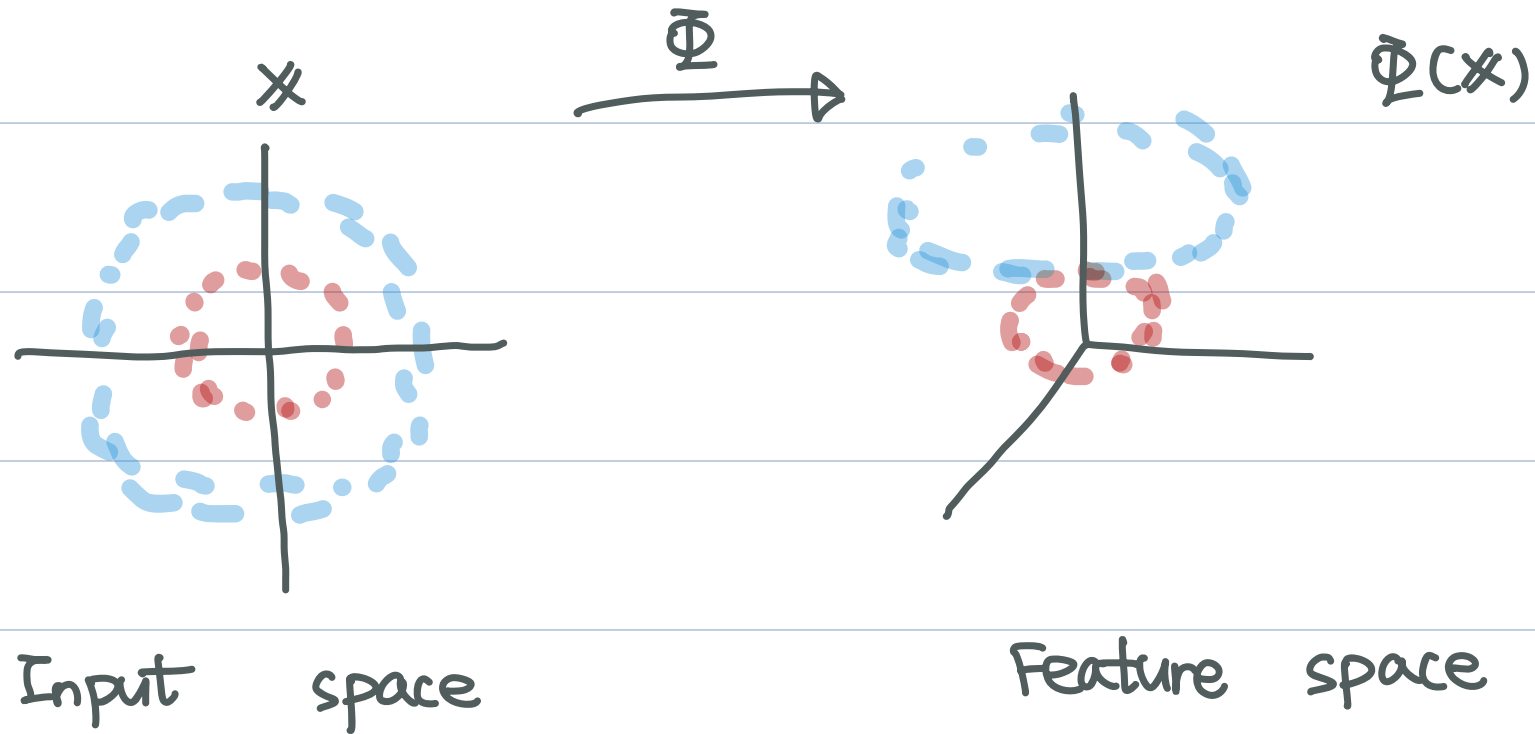
Kernel SVM

kernel method (kernel trick or technique)

Purpose : 예측, 분류 모델에 non-linearity를 부여
(Φ , w 모두)

non-linearity 를 주는 다른 방법?

- feature transformation or basis function



공간 변환

$x_1, x_2 \dots x_N$

$\Phi(x_1), \Phi(x_2) \dots \Phi(x_N)$

Kernel method

방법 (key idea)

- If an algorithm only uses dot products of vectors, then we can replace the dot product with a kernel function

$$\mathbf{x}^T \mathbf{x}'$$

$$\longrightarrow K(\mathbf{x}, \mathbf{x}')$$

$$\Phi(\mathbf{x})^T \Phi(\mathbf{x}')$$

유사도 또는 상관 정도 측정

How to construct a valid kernel function?

Generalized linear regression

Linear model : $y(x) = w^T x + b$

Basis function method

1. Fix a basis function $\Phi(x) = (\phi_0(x), \phi_1(x), \dots, \phi_{M-1}(x))^T$

2. Model $y(x) = w^T \Phi(x)$ with weight vector w

3. Finding w minimizing the cost function below

$$J(w) = \frac{1}{2} \sum_{n=1}^N \{ w^T \Phi(x_n) - t_n \}^2 + \frac{\lambda}{2} w^T w$$

where $\lambda \geq 0$.

4. Inference : $y(\hat{x}) = w^T \Phi(\hat{x})$

$$J(w) = \frac{1}{2} (\Phi w - t)^T (\Phi w - t) + \frac{\lambda}{2} w^T w$$

Set $\nabla_w J(w) = 0$. Then we see

$$w = \sum_{n=1}^N a_n \Phi(x_n) = \underbrace{\Phi^T}_{M \times N} a^{N \times 1} \quad (6.3)$$

where Φ is the design matrix and $a = (a_1, \dots, a_N)^T$ with

$$a_n = -\frac{1}{\lambda} \{ w^T \Phi(x_n) - t_n \}$$

Substitute $w_1 = \Phi^T \alpha$ into $J(w_1)$ (6.2)

$$J(\alpha) = \frac{1}{2} \alpha^T \Phi \Phi^T \Phi \Phi^T \alpha - \alpha^T \Phi \Phi^T \mathbf{t} + \frac{1}{2} \mathbf{t}^T \mathbf{t} + \frac{\lambda}{2} \alpha^T \Phi \Phi^T \alpha$$

where $\mathbf{t} = (t_1, t_2, \dots, t_N)^T$.

$$\Phi \Phi^T = \begin{pmatrix} -\Phi(x_1)^T - \\ -\Phi(x_2)^T - \\ \vdots \\ -\Phi(x_N)^T - \end{pmatrix} \begin{pmatrix} 1 & 1 & \dots & 1 \\ \Phi(x_1) & \Phi(x_2) & \dots & \Phi(x_N) \\ 1 & 1 & \dots & 1 \end{pmatrix} \quad N \times N$$

$$(\Phi \Phi^T)_{ij} = \Phi(x_i)^T \Phi(x_j) \quad \searrow \text{replace } K \quad N \times N$$

$$K_{ij} = K(x_i, x_j)$$

kernel trick (dual representation)

1. Fix a kernel function $k(x, x')$

2. Model $y(x) = \mathbf{K}(x)^T \alpha$ where $\mathbf{K}(x) = (k_1(x), k_2(x), \dots, k_N(x))^T$

and $k_i(x) = k(x_i, x)$

3. Finding α minimizing the cost function below

$$J(\alpha) = \frac{1}{2} \alpha^T K K \alpha - \alpha^T K \mathbf{t} + \frac{1}{2} \mathbf{t}^T \mathbf{t} + \frac{\gamma}{2} \alpha^T K \alpha$$

where $k_{nm} = k(x_n, x_m)$, $\gamma \geq 0$

4. Inference $y(\hat{x}) = \mathbf{K}(\hat{x})^T \alpha = \sum_{i=1}^N \alpha_i k(x_i, \hat{x})$

kernel support vector machine

Primal problem

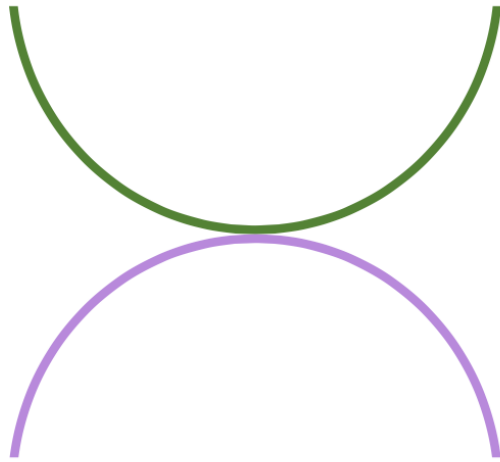
$$\min_{w, b} \frac{1}{2} \|w\|^2 \quad \text{s.t.} \quad t_{\tilde{\lambda}} (w^T \tilde{x}_{\tilde{\lambda}} + b) \geq 1 \quad \forall \tilde{\lambda} = 1, \dots, N$$

$$\tilde{x}^T \tilde{x}'$$

Primal problem vs Dual problem

Primal problem : original objective function

minimize objective function



Dual problem : lower bound

maximize lower bound

Primal problem (hard margin)

$$\min_{w, b} \frac{1}{2} \|w\|^2 \quad \text{s.t.} \quad t_i (w^T x_i + b) \geq 1 \quad \forall i = 1, \dots, N$$

Construct Lagrangian

To handle the inequality, introduce Lagrangian multipliers

$\alpha_i \geq 0$ and form the Lagrangian:

$$\mathcal{L}(w, b, \alpha) = \frac{1}{2} \|w\|^2 - \sum_{i=1}^N \alpha_i (t_i (w^T x_i + b) - 1)$$

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

$$\alpha := (\alpha_1, \dots, \alpha_N)^T$$

Derive the dual

To obtain the dual, we minimize the Lagrangian w.r.t. w and b , and maximize w.r.t. α . I.e.

$$\min_{w, b} \max_{\alpha \geq 0} \mathcal{L}(w, b, \alpha)$$

$$-\frac{\partial \mathcal{L}}{\partial w} = 0 \quad \Rightarrow \quad w = \sum_{i=1}^N \alpha_i t_i x_i$$

$$-\frac{\partial \mathcal{L}}{\partial b} = 0 \quad \Rightarrow \quad \sum_{i=1}^N \alpha_i t_i = 0$$

Substitute these back into the Lagrangian

Eliminating w and b from $L(w, b, x)$ using these conditions. Then this gives the dual representation

$$\tilde{L}(x) := \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j} \alpha_i \alpha_j t_i t_j \underbrace{x_i^T x_j}_{\text{dot product}}$$

$$\text{s.t.} \quad \alpha_i \geq 0 \quad i=1, \dots, N, \quad \sum_{i=1}^N \alpha_i t_i = 0$$

Dual optimization problem

$$\max_{\alpha} \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j} \alpha_i \alpha_j t_i t_j \mathbf{x}_i^T \mathbf{x}_j$$

$$\text{s.t.} \quad \alpha_i \geq 0 \quad i=1, \dots, N, \quad \sum_{i=1}^N \alpha_i t_i = 0$$

Kernel trick (Kernel SVM)

If the data is not linearly separable, we map the input data to a higher-dimensional space using a basis function $\Phi(\cdot)$ or we replace the dot product $\mathbf{x}^T \mathbf{x}'$ with a kernel $K(\mathbf{x}, \mathbf{x}')$. I.e.

$$\max_{\alpha} \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j} \alpha_i \alpha_j t_i t_j K(\mathbf{x}_i, \mathbf{x}_j)$$

$$\text{s.t.} \quad \alpha_i \geq 0 \quad i = 1, \dots, N, \quad \sum_{i=1}^N \alpha_i t_i = 0$$

Prediction function

Once the optimal α_i are found, the decision function for a new input x is :

$$f(x) = \sum_{i=1}^N \alpha_i t_i K(x_i, x) + b$$

where b is a bias.

SVM: Overlapping classes (recall)

$$\min_{w, \xi} \frac{1}{2} \|w\|^2 + c \cdot \sum_{i=1}^N \xi_i$$

primal problem

$$\text{s.t.} \quad t_i (w^T x_i + b) \geq 1 - \xi_i$$

$$\forall i = 1, \dots, N$$

$$\xi_i \geq 0$$

Here $c > 0$ controls the trade-off between the slack variable penalty and the margin

To obtain the dual, we minimize the Lagrangian w.r.t. w and b , and maximize w.r.t. α and μ

$$\min_{w, b} \max_{\alpha, \mu \geq 0} \mathcal{L}(w, b, \xi, \alpha, \mu)$$

where $\mathcal{L}(w, b, \xi, \alpha, \mu)$

$$= \frac{1}{2} \|w\|^2 + C \sum_{i=1}^N \xi_i - \sum_{i=1}^N \alpha_i (t_i \gamma(x_i) - 1 + \xi_i) - \sum_{i=1}^N \mu_i \xi_i$$

$\{\alpha_i \geq 0\}$ and $\{\mu_i \geq 0\}$ are Lagrange multipliers.

$$\alpha_i (t_i \gamma(x_i) - 1 + \xi_i)$$

We optimize out w , b and $\{\xi_i\}$

$$-\frac{\partial \mathcal{L}}{\partial w} = 0 \Rightarrow w = \sum_{i=1}^N \alpha_i t_i x_i$$

$$-\frac{\partial \mathcal{L}}{\partial b} = 0 \Rightarrow \sum_{i=1}^N \alpha_i t_i = 0$$

$$-\frac{\partial \mathcal{L}}{\partial \xi_i} = 0 \Rightarrow \alpha_i = C - \mu_i$$

Using these results to eliminate w , b and $\{\xi_i\}$ from the Lagrangian, we obtain the dual Lagrangian

$$\tilde{\mathcal{L}}(\alpha) = \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j} \alpha_i \alpha_j t_i t_j x_i^T x_j$$

We note that $\alpha_i \geq 0$ and $\alpha_i = C - \mu_i$

Since $\mu_i \geq 0$, we have $\alpha_i \leq C$

Therefore we have to minimize $\tilde{L}(\alpha)$ w.r.t $\{\alpha_i\}$

subject to

$$0 \leq \alpha_i \leq C \quad \text{for } i=1, \dots, N$$

$$\sum_{i=1}^N \alpha_i t_i = 0$$

Kernel SVM: Overlapping classes

$$\max_{\alpha} \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j} \alpha_i \alpha_j t_i t_j K(x_i, x_j)$$

$$\text{s.t.} \quad 0 \leq \alpha_i \leq C \quad \forall i=1, \dots, N, \quad \sum_{i=1}^N \alpha_i t_i = 0 //$$