

Kernels

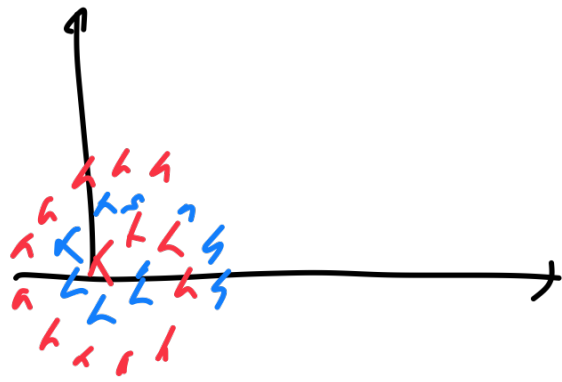
let us consider a long feature vector

$$\tilde{\phi}(\tilde{x}^{(i)}) = (1, x^{(i)}, (x^{(i)})^2, (x^{(i)})^3, \dots, (x^{(i)})^d)$$

We want to minimize the loss

$$l(\beta) = \frac{1}{N} \sum_{i=1}^N (t^{(i)} - \beta^T \tilde{\phi}(x^{(i)}))^2$$

$$\beta^{(k+1)} = \beta^{(k)} + \eta \frac{1}{N} \sum_{i=1}^N (t^{(i)} - \beta^T \tilde{\phi}(x^{(i)})) \cdot \tilde{\phi}(x^{(i)})$$



$$X = (x_1, \dots, x_D)$$

$$D = 1000$$

$$\tilde{\Phi} = 10^9$$

$$\tilde{\Phi}(x) = \left[\begin{array}{c} 1 \\ x_1 \\ \vdots \\ x_D \\ x_1^2 \\ \vdots \\ x_D^2 \\ x_1^3 \\ \vdots \\ x_1 x_2 x_3 \\ \vdots \end{array} \right] \rightarrow 1 + D + D^2 + D^3$$

Idea: Write $\beta^{(k)}$ as $\beta^{(k)} = \sum_{j=1}^N \Phi(x^{(j)}) \lambda_j$

Start $\beta^{(0)} = \vec{0}$

$$\beta^{(k+1)} = \beta^{(k)} + \frac{\eta}{N} \sum_{i=1}^N (t^{(i)} - \beta^T \tilde{\phi}(x^{(i)})) \cdot \tilde{\phi}^{(i)}$$

$$= \sum_{j=1}^N \underbrace{\phi(x^{(j)})}_{\downarrow} \lambda_j + \frac{\eta}{N} \sum_{i=1}^N (t^{(i)} - \sum_{j=1}^N \lambda_j \phi(x^{(j)})^T \underbrace{\phi(x^{(i)})}_{\tilde{\phi}^{(i)}})$$

$$= \sum_{j=1}^N \underbrace{\phi(x^{(j)})}_{\downarrow} (\lambda_j + \frac{\eta}{N} (t^{(j)} - \sum_{i=1}^N \lambda_i \phi(x^{(i)})^T \phi(x^{(j)})))$$

$$\lambda_j^{(k+1)} = \lambda_j^{(k)} + \frac{\eta}{N} (t^{(j)} - \sum_{i=1}^N \lambda_i \phi(x^{(i)})^T \phi(x^{(j)}))$$

Kernels

$$y(x) = \beta^T \phi(x) \quad \leftarrow \quad \sum \lambda_i^{(k)} \phi(x^{(i)})$$

let $\phi(x^{(i)})$ denote a feature vector (e.g. polynomial features)

$$\beta^{(k+1)} \leftarrow \beta^{(k)} + \eta \sum_{i=1}^N (\underbrace{t^{(i)}} - \underbrace{\beta^T \phi(x^{(i)})}) \cdot \phi(x^{(i)})$$

$\phi(x^{(i)})$ depends on D (dimension) $\rightarrow$ can be expensive to store when D is large

How about storing β as $\sum_{i=1}^N \lambda_i \phi(x^{(i)})$

$$\sum_{i=1}^N \lambda_i^{(0)} \phi(x^{(i)})$$

Always possible to use such a decomposition for $\beta^{(0)} = \mathbf{0}$

$$\begin{aligned}\beta^{(k+1)} &= \sum_{i=1}^N \lambda_i \phi(x^{(i)}) + \eta \sum_{i=1}^N (t^{(i)} - \sum_{j=1}^N \lambda_j \phi(x^{(j)})^T \phi(x^{(i)})) \phi(x^{(i)}) \\ &= \sum_{i=1}^N \phi(x^{(i)}) \left(\lambda_i + \eta \left(t^{(i)} - \sum_{j=1}^N \lambda_j \phi(x^{(j)})^T \phi(x^{(i)}) \right) \right)\end{aligned}$$

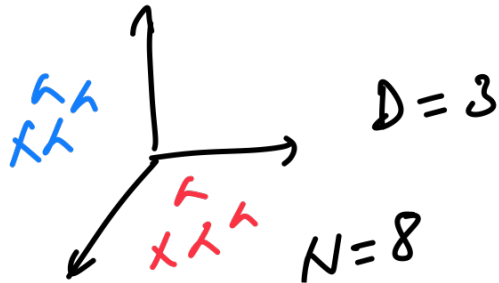
$$\beta^{(k)} = \sum_{i=1}^N \lambda_i^{(k)} \phi(x^{(i)}) \rightarrow \lambda^{(k+1)}$$

$$\beta^{(k+1)} = \sum_{i=1}^N \lambda_i^{(k+1)} \phi(x^{(i)})$$

$$\lambda_i^{(k+1)} = \lambda_i^{(k)} + \eta \left(t^{(i)} - \sum_{j=1}^N \lambda_j^{(k)} \underbrace{\phi(x^{(j)})^T \phi(x^{(i)})}_{\uparrow} \right) \quad (*)$$

General idea: iterate over the $\lambda_i^{(k)}$ instead $\beta_j^{(k)}$

→ storage goes from $O(D)$ to N



$$K(x^{(i)}, x^{(j)}) = K(i, j) = \phi(x^{(i)})^T \phi(x^{(j)})$$

→ K is of size N^2

Why is it interesting to have $\langle \phi(x^{(i)}), \phi(x^{(j)}) \rangle$
 $= \phi(x^{(i)})^T \phi(x^{(j)})$

Consider polynomial features of max degree 3

$$\phi(x) = [1, x_1, x_2, x_3, \dots, x_1^2, x_2^2, \dots, x_1^3, x_2^3, \dots, x_D^3]$$

$$\phi(z) = [1, z_1, z_2, z_3, \dots, z_1^2, z_2^2, \dots, z_1^3, z_2^3, \dots, z_D^3]$$

$$\begin{aligned} \phi(x)^T \phi(z) &= 1 + \sum x_i z_i + \sum x_i x_j z_j z_j + \sum x_i x_j x_k z_i z_j z_k \\ &= 1 + (x^T z) + (x^T z)^2 + (x^T z)^3 \end{aligned}$$

$$\lambda_i^{(k+1)} \leftarrow \lambda_i^{(k)} + \eta (t^{(i)} - \sum_{j=1}^N \lambda_j^{(k)} \phi(x^{(i)})^T \phi(x^{(j)}))$$

$$\leftarrow \lambda_i^{(k)} + \eta (t^{(i)} - \sum_{j=1}^N \lambda_j^{(k)} K(i, j))$$

If we want to focus on K instead of ϕ
 What are the "good" K matrices that we could use
 to learn the properties of our data?

$K \rightarrow$ Square Symmetric

$\rightarrow$ positive semidefinite (psd)

To see why K has to be psd

take $z \in \mathbb{R}^N$

$$\text{if } \exists \phi(x^{(i)}), \phi(x^{(j)}) \text{ s.t. } K(i, j) \\ = \phi(x^{(i)})^T \phi(x^{(j)})$$

then

$$\underline{z^T K z} = \sum_{i,j} z_i K_{ij} z_j$$

$$= \sum_{i,j} z_i \phi(x^{(i)})^T \phi(x^{(j)}) z_j$$

$$= \sum_{i,j,k} z_i \phi_k(x^{(i)}) \phi_k(x^{(j)}) z_j$$

$$= \sum_k \left(\sum_i \underline{z_i \phi_k(x^{(i)})} \right) \left(\sum_j \underline{z_j \phi_k(x^{(j)})} \right)$$

$$= \sum_k \left(\sum_i z_i \phi_k(x^{(i)}) \right)^2$$

$\Rightarrow z^T K z \geq 0 \Leftrightarrow K$ positive semidefinite matrix

$\Rightarrow$ Theorem (Mercer) For k to be a valid kernel (Mercer kernel) it is necessary and sufficient to have K symmetric and positive semidefinite

$$K(i,j) = \frac{1}{\sigma} \exp \left(-\frac{\|x^{(i)} - x^{(j)}\|^2}{\sigma^2} \right)$$

λ_i we get after
our
iterations

How do we recover $y(x)$ from λ_i^* ?

$$\begin{aligned} y(x) &= \beta^T \phi(x) = \left(\sum_{i=1}^N \lambda_i^* \phi(x^{(i)}) \right)^T \phi(x) \\ &= \sum_{i=1}^N \lambda_i^* \phi(x^{(i)})^T \phi(x) \\ &= \lambda^T \underbrace{K(x^{(i)}, x)} \end{aligned}$$

as an example if $K(i,j)$ can be generated from a
function $k(x,y)$ of the original data,

We can get the expression of our model $y(x)$ directly from that function $k(x, y)$

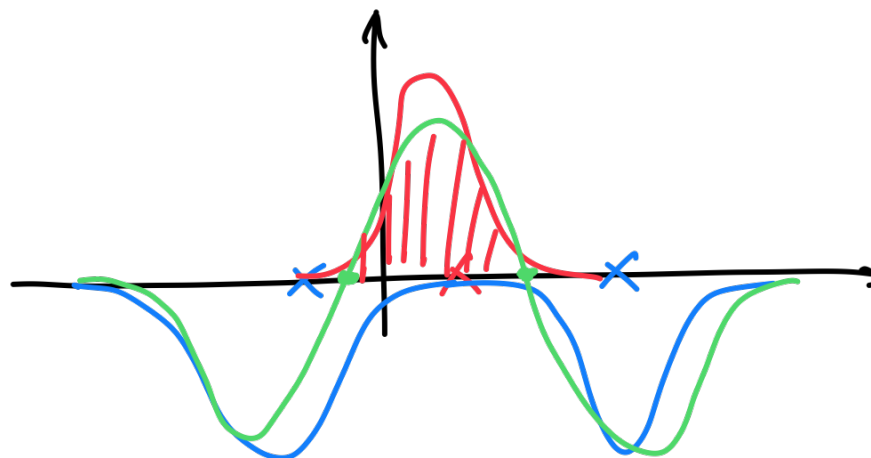
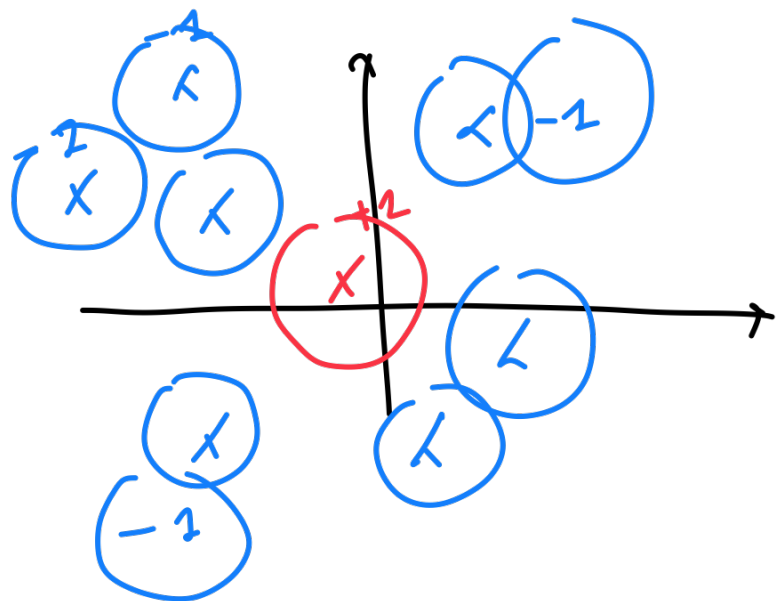
e.g take $k(\underline{x}, \underline{y}) = \frac{1}{\sigma} \exp\left(-\frac{\|\underline{x} - \underline{y}\|^2}{\sigma^2}\right)$

To learn $y(x)$ we first need the N^2

$$K(i, j) = \frac{1}{\sigma} \exp\left(-\frac{\|x^{(i)} - x^{(j)}\|^2}{\sigma^2}\right) \text{ for every } x^{(i)}, x^{(j)}$$

then learn the $\lambda_i^{(k)}$

Finally you get your model as $y(x) = \sum_{i=1}^N \lambda_i^* \exp\left(-\frac{\|x^{(i)} - x\|^2}{\sigma^2}\right)$



Kernels

General idea $\rightarrow$ Switch from an approach based on feature vectors (here on the dimension D)

to an approach depending on the number of training points N

$$\rightarrow \mathcal{L}(\beta) = \frac{1}{N} \sum_{i=1}^N (t^{(i)} - \beta^T \phi(x^{(i)}))^2$$

$$\beta^{(k+1)} \leftarrow \beta^{(k)} + \eta \frac{1}{N} \sum_{i=1}^N (t^{(i)} - \beta^T \phi(x^{(i)})) \phi(x^{(i)})$$

general idea $\beta^{(k)} = \sum \lambda_i^{(k)} \phi(x^{(i)})$

$$\begin{aligned} \beta^{(k+1)} &\leftarrow \sum_{i=1}^N \lambda_i^{(k)} \phi(x^{(i)}) + \frac{\eta}{N} \sum_{i=1}^N (t^{(i)} - \sum_{j=1}^N \lambda_j^{(k)} \phi(x^{(j)})^T \phi(x^{(i)})) \phi(x^{(i)}) \\ &\leftarrow \sum_{i=1}^N \left(\lambda_i^{(k)} + \frac{\eta}{N} \sum_{i=1}^N (t^{(i)} - \sum_{j=1}^N \lambda_j^{(k)} \phi(x^{(j)})^T \phi(x^{(i)})) \right) \phi(x^{(i)}) \end{aligned}$$

$$\lambda_i^{(k+1)} \leftarrow \lambda_i^{(k)} + \frac{\eta}{N} \left(t^{(i)} - \sum_{j=1}^N \lambda_j^{(k)} \phi(x^{(j)})^T \phi(x^{(i)}) \right)$$

$\underbrace{\phi(x^{(j)})^T \phi(x^{(i)})}_{K(i,j)} \lambda_j$

We then introduce $K(x^{(i)}, x^{(j)})$ to store the scalar products between the feature vectors.

We can then study the class of K matrices for which there exist a decomposition in feature vectors

→ Mercer's kernel : positive semi-definite K
(symmetric)

e.g $K(x, y) = \exp\left(-\frac{\|x - y\|^2}{\sigma}\right)$

Claim $K(x, y) = \exp\left(-\frac{\|x - y\|^2}{\sigma}\right)$ is a valid kernel
 (meaning $\exists \phi(x^{(i)}) \phi(x^{(j)})$)

Such that

$$K(x^{(i)}, x^{(j)}) = \phi(x^{(i)})^T \phi(x^{(j)})$$

Proof

$$\exp\left(-\frac{\|x^{(i)} - x^{(j)}\|^2}{\sigma}\right) = e^{-\frac{1}{\sigma} \|x^{(i)}\|^2} e^{-\frac{1}{\sigma} \|x^{(j)}\|^2} e^{\frac{2}{\sigma} x^{(i)T} x^{(j)}}$$

$$= e^{-\frac{1}{\sigma} \|x^{(i)}\|^2} e^{-\frac{1}{\sigma} \|x^{(j)}\|^2}$$

Taylor

$$\left(1 + \frac{2}{\sigma} \sum_{k=1}^D x_k^{(i)} x_k^{(j)} + \left(\frac{2}{\sigma}\right)^2 \frac{\left(\sum_k x_k^{(i)} x_k^{(j)}\right)^2}{2!} \right. \\ \left. + \frac{1}{3!} \left(\frac{2}{\sigma}\right)^3 \left(\sum_k x_k^{(i)} x_k^{(j)}\right)^3 + \dots \right)$$

$$\exp\left(-\frac{\|x^{(i)} - x^{(j)}\|^2}{\sigma}\right) = e^{-\frac{1}{\sigma} \|x^{(i)}\|^2} e^{-\frac{1}{\sigma} \|x^{(j)}\|^2} \quad \text{cloud} = \phi(x^{(i)})^T \phi(x^{(j)})$$

$$* \left(1, \sqrt{\frac{2}{\sigma}} x^{(i)}, \sqrt{\left(\frac{2}{\sigma}\right)^2 \frac{1}{2!}} (x_k^{(i)} x_l^{(i)})_{k,l}, \right. \\ \left. \sqrt{\left(\frac{2}{\sigma}\right)^3 \frac{1}{3!}} (x_k^{(i)} x_l^{(i)} x_m^{(i)})_{k,l,m}, \dots \right)^T$$

$$\left(\sum_{k=1}^D x_k^{(i)} x_k^{(j)} \right)^2$$

$$= \sum_{k=1}^D \sum_{l=1}^D x_k^{(i)} x_l^{(i)} x_k^{(j)} x_l^{(j)} =$$

$$\left(1, \sqrt{\frac{2}{\sigma}} x^{(j)}, \sqrt{\left(\frac{2}{\sigma}\right)^2 \frac{1}{2!}} (x_k^{(j)} x_l^{(j)})_{k,l}, \right. \\ \left. \dots \right)$$

Formal statement of Mercer:

X measure space, We will $k: L^2(X \times X) \rightarrow \mathbb{R}$ is a valid kernel iff there exist some feature map $\varphi: X \rightarrow H$ H separable Hilbert space

$$k(x, x') = \langle \varphi(x), \varphi(x') \rangle_H$$

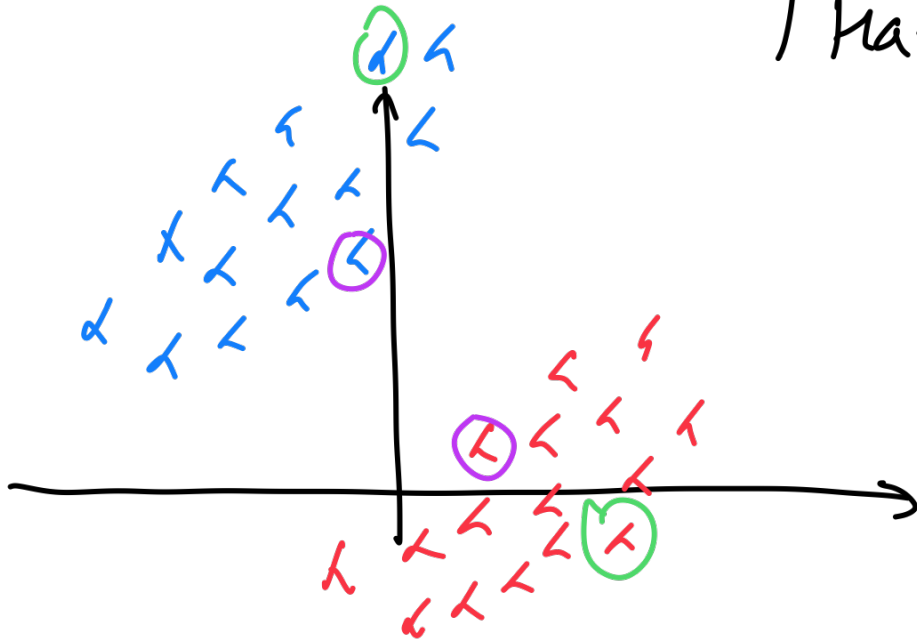
Given a kernel k we say that k satisfies Mercer's condition if

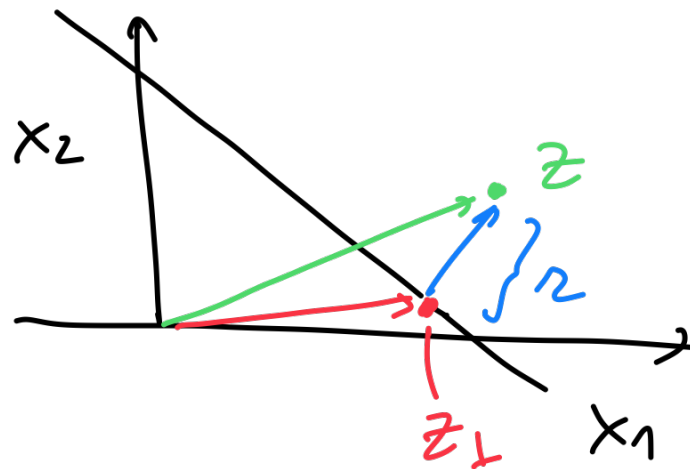
$$\int_{X \times X} k(x, x') f(x) f(x') dx dx' \geq 0 \quad \text{for all } f \in L^2(X)$$

Merger theorem

k valid kernel iff k is Mercer

Support Vector Machines / Kernel Vector Machines
/ Sparse Vector Machines
/ Maximum Margin Classifier

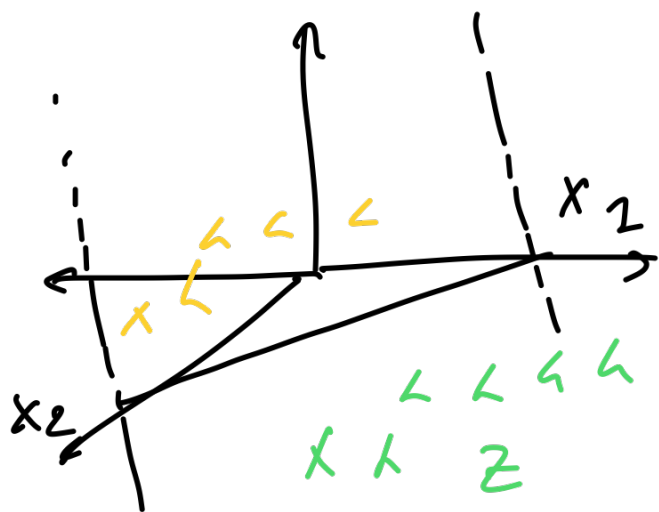




$$z = z_{\perp} + \frac{(\beta_1, \beta_2)}{\|(\beta_1, \beta_2)\|} \cdot \underbrace{\quad}_{\text{norm vector}} \quad ?$$

plane $\beta_1 x_1 + \beta_2 x_2 + \beta_0 = 0$

normal vector $= (\beta_1, \beta_2)$



$$x_3 = \beta_1 x_1 + \beta_2 x_2 + \beta_0$$

$$l(\beta) = \frac{1}{N} \sum_{i=1}^N (t^{(i)} - (\beta_0 + \beta_1 x_1 + \beta_2 x_2))^2$$

$$x_2 = -\frac{\beta_1}{\beta_2} x_1 - \frac{\beta_0}{\beta_2}$$

$$\underset{\substack{\parallel \\ (z_1, z_2)}}{z} = z_{\perp} + \frac{(\beta_1, \beta_2)}{\|(\beta_1, \beta_2)\|} \cdot r = ((z_{\perp})_1, (z_{\perp})_2) + \frac{(\beta_1, \beta_2)}{\|(\beta_1, \beta_2)\|} r$$

$$\beta_0 + \beta_1 (z_{\perp})_1 + \beta_2 (z_{\perp})_2 = 0$$

$$\beta_0 + \beta_1 z_1 + \beta_2 z_2 = \beta_0 + \beta_1 (z_{\perp})_1 + \beta_2 (z_{\perp})_2 \overset{=0}{\quad}$$

$$\underbrace{\hspace{1.5cm}}_{y(z)}$$

$$+ \frac{\beta_1 \beta_1}{\|(\beta_1, \beta_2)\|} r + \frac{\beta_2 \beta_2}{\|(\beta_1, \beta_2)\|} r$$

$$= r \frac{\|(\beta_1, \beta_2)\|^2}{\|(\beta_1, \beta_2)\|} = r (\| \beta_1, \beta_2 \|)$$

For any point z , the distance r of z to the plane is
given by

$$r = \frac{y(z)}{\|(\beta_1, \beta_2)\|}$$

$$\text{dist} = r t(z) = \frac{y(z) \cdot t(z)}{\|(\beta_1, \beta_2)\|}$$

From r we can
get $|r|$ by

multiplying by
 $t^{(i)}$

(target of $z^{(i)}$)

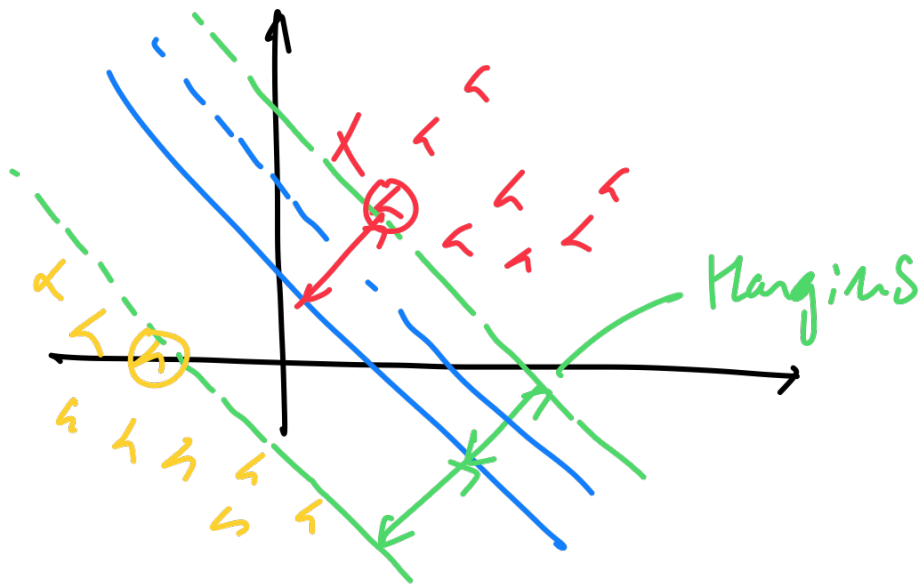
provided $t^{(i)} = +1$ if
 $z^{(i)}$ is
above

-1 if $z^{(i)}$ is
below

We can then look for the plane that maximizes its distance with respect to the closest $z^{(i)}$

$$\max_{(\beta_0, \beta_1, \beta_2)} \text{ with } z^{(i)} \frac{t(z^{(i)})y(z^{(i)})}{\|(\beta_1, \beta_2)\|} = \frac{t(z^{(i)}) (\beta_0 + \beta_1 z_1^{(i)} + \beta_2 z_2^{(i)})}{\|(\beta_1, \beta_2)\|}$$

$\xrightarrow{\pm 1}$



ratio is
invariant by
rescaling of β

$$\beta \leftarrow \alpha \beta$$

$$\frac{t(z^{(i)}) (\cancel{\beta_0} + \cancel{\beta_1} z_1^{(i)} + \cancel{\beta_2} z_2^{(i)})}{\cancel{\beta_1} \parallel (\beta_1, \beta_2) \parallel} \rightarrow$$

→ let us choose β such that for the closest $z^{(i)}$ we have

$$t(z^{(i)}) (\beta_0 + \beta_1 z_1^{(i)} + \beta_2 z_2^{(i)}) = 1$$

⇒ For the closest point the distance is $\frac{1}{\parallel (\beta_1, \beta_2) \parallel}$)

⇒ For all the other points, necessarily we have

$$\text{distance} \geq \frac{1}{\parallel (\beta_1, \beta_2) \parallel}$$

For all the other points $t(z^{(j)}) (\beta_0 + \beta_1 z_1^{(j)} + \beta_2 z_2^{(j)}) \geq 1$

the problem then turns into

$$\max_{(\beta_0, \beta_1, \beta_2)} \frac{1}{\|(\beta_1, \beta_2)\|}$$

under the constraints

$$t^{(i)} (\beta_0 + \beta_1 z_1^{(i)} + \beta_2 z_2^{(i)}) \geq 1$$

for every i

We can further simplify the formulation into

$$\min_{\beta} \quad \|(\beta_1, \beta_2)\|^2$$

β

$$t^{(i)} (\beta_0 + \beta_1 z_1^{(i)} + \beta_2 z_2^{(i)}) \geq 1$$