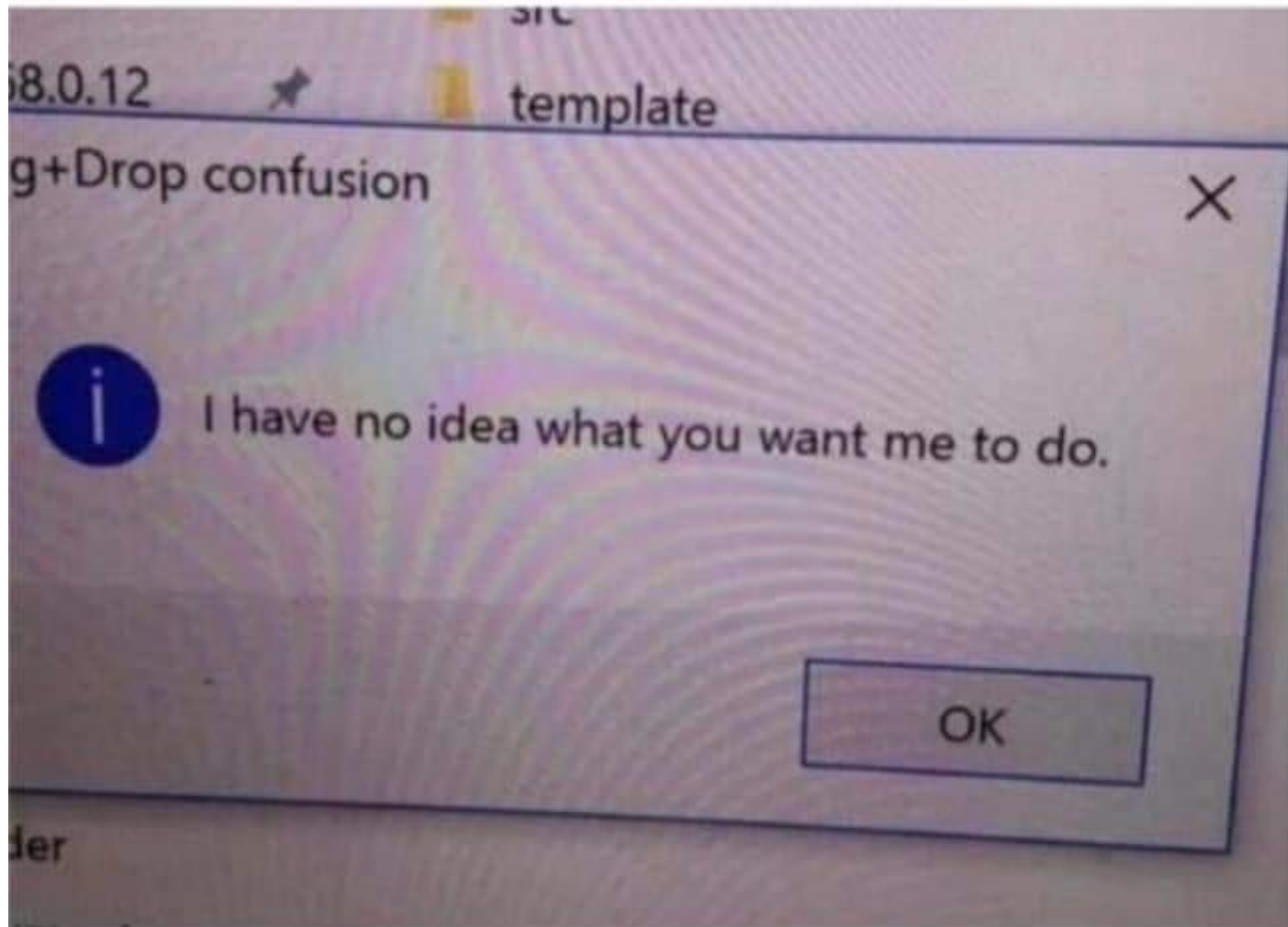


Kernel SVM

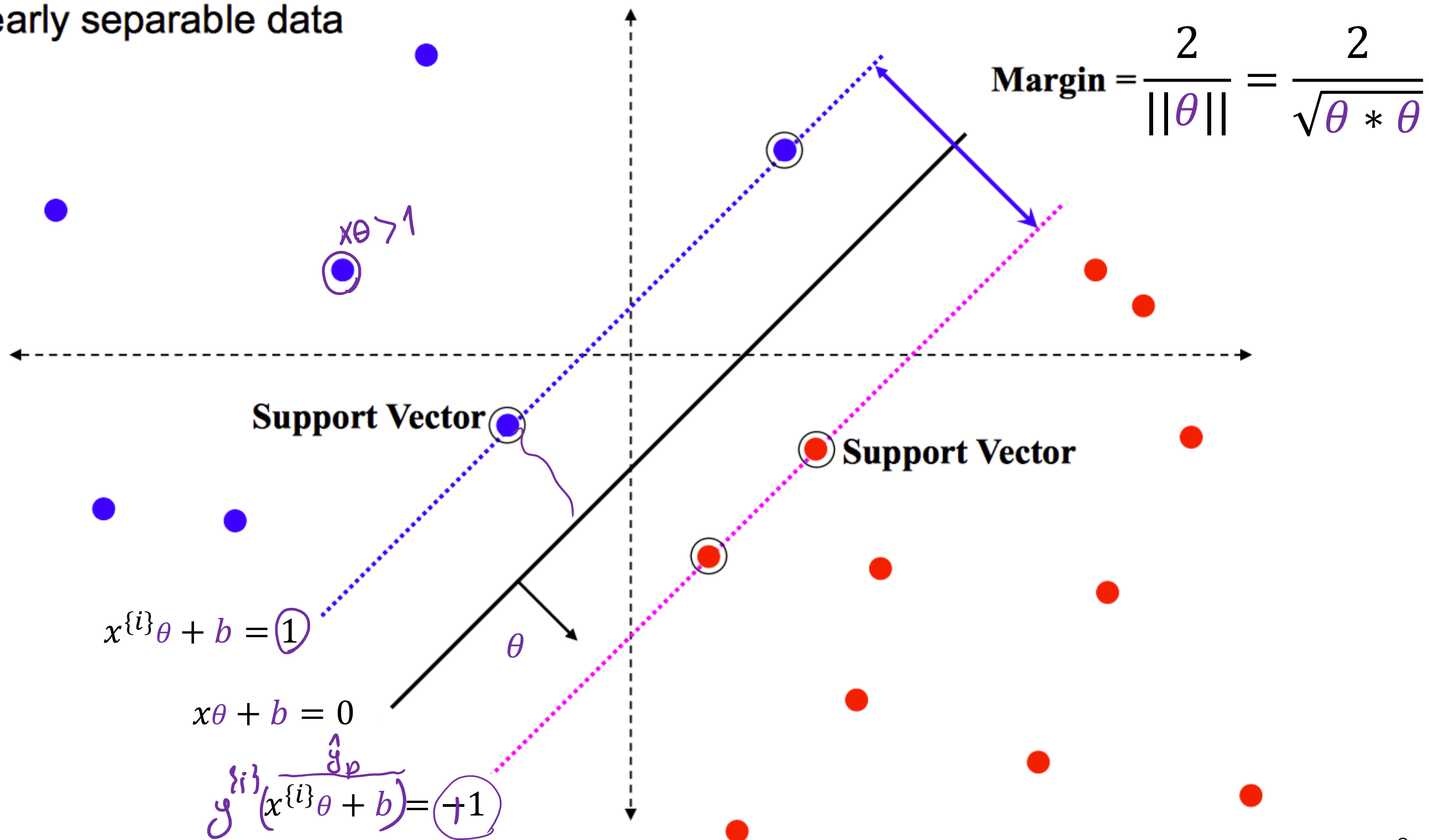
Mahdi Roozbahani
Georgia Tech

When you get angry at the computer
because your code doesn't work

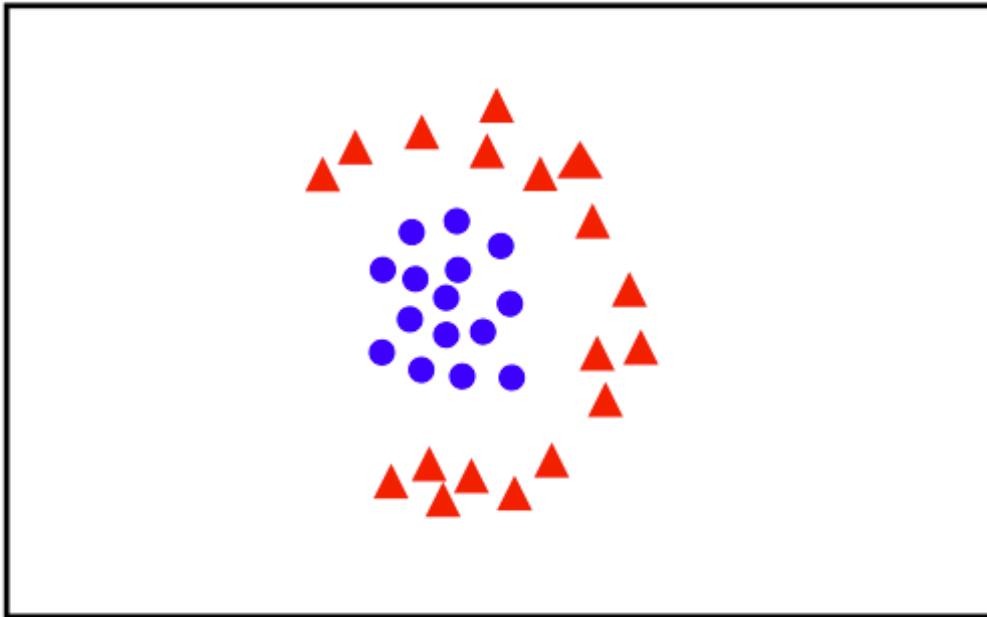


Geometric Interpretation

linearly separable data



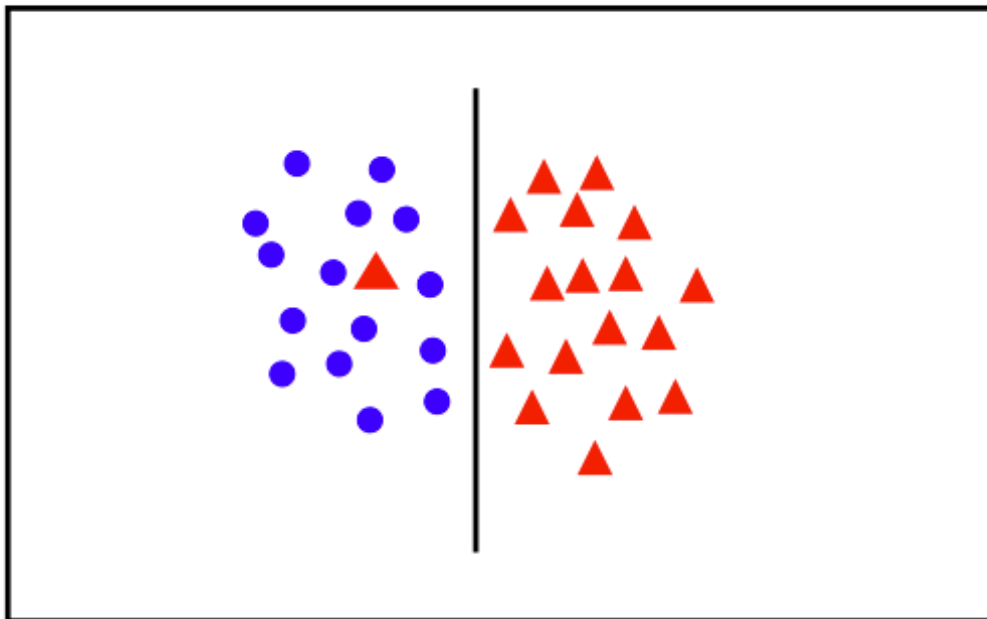
Handling Data that are Not Linearly Separable



- linear classifier not appropriate

??

Kernel trick

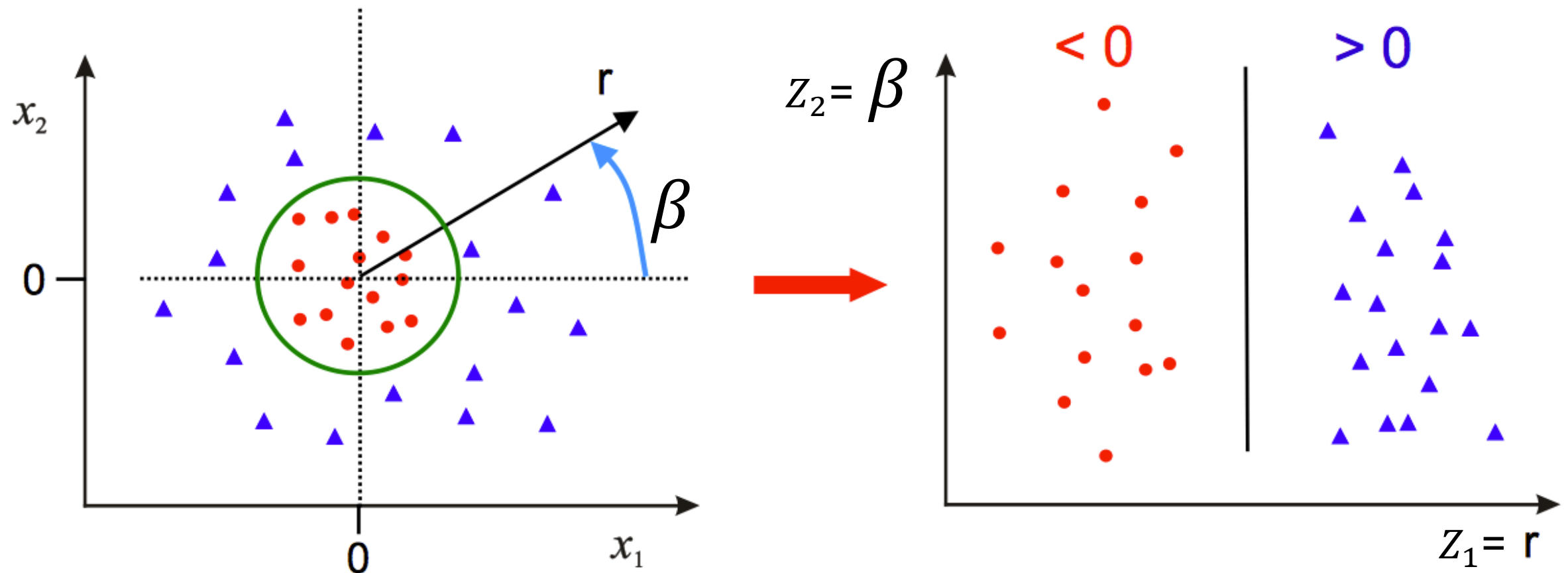


- introduce slack variables

Soft Margin SVM

(allowing ourselves to make errors)

Idea 1: Use Polar Coordinates to go to z space

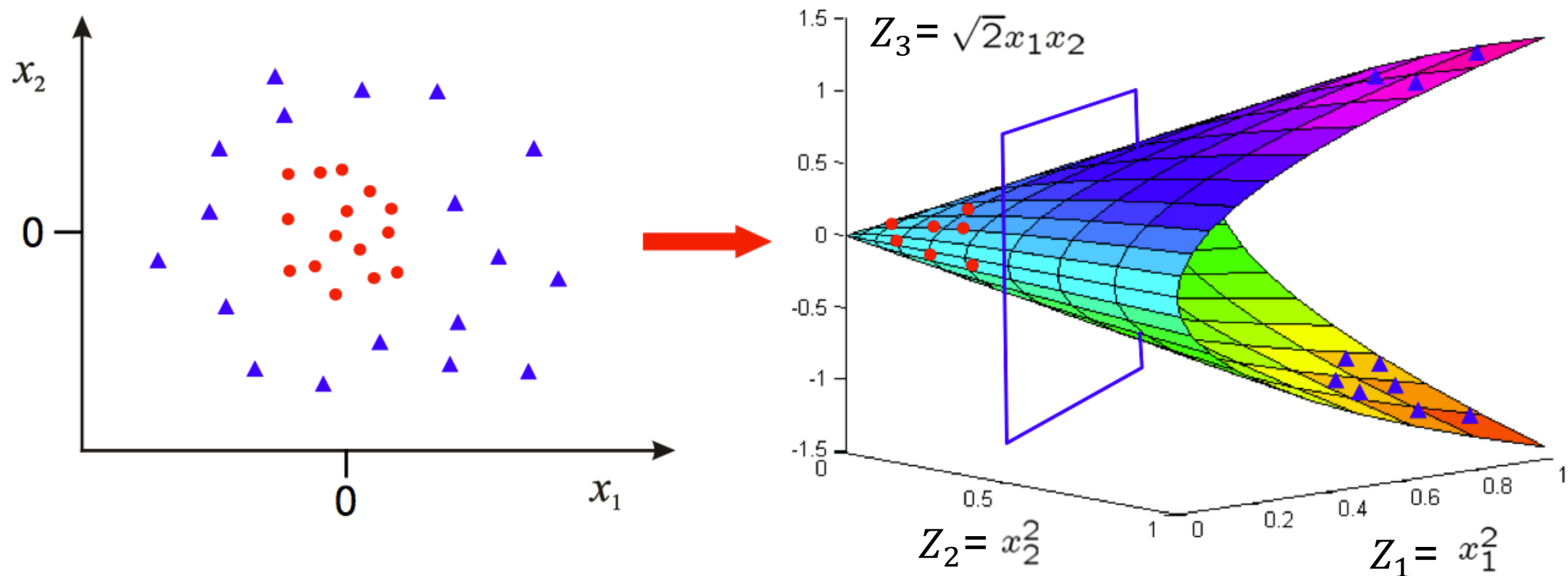


- Data **is** linearly separable in polar coordinates
- Acts non-linearly in original space

$$\Phi : \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \rightarrow \begin{pmatrix} r \\ \beta \end{pmatrix} \quad \mathbb{R}^2 \rightarrow \mathbb{R}^2$$

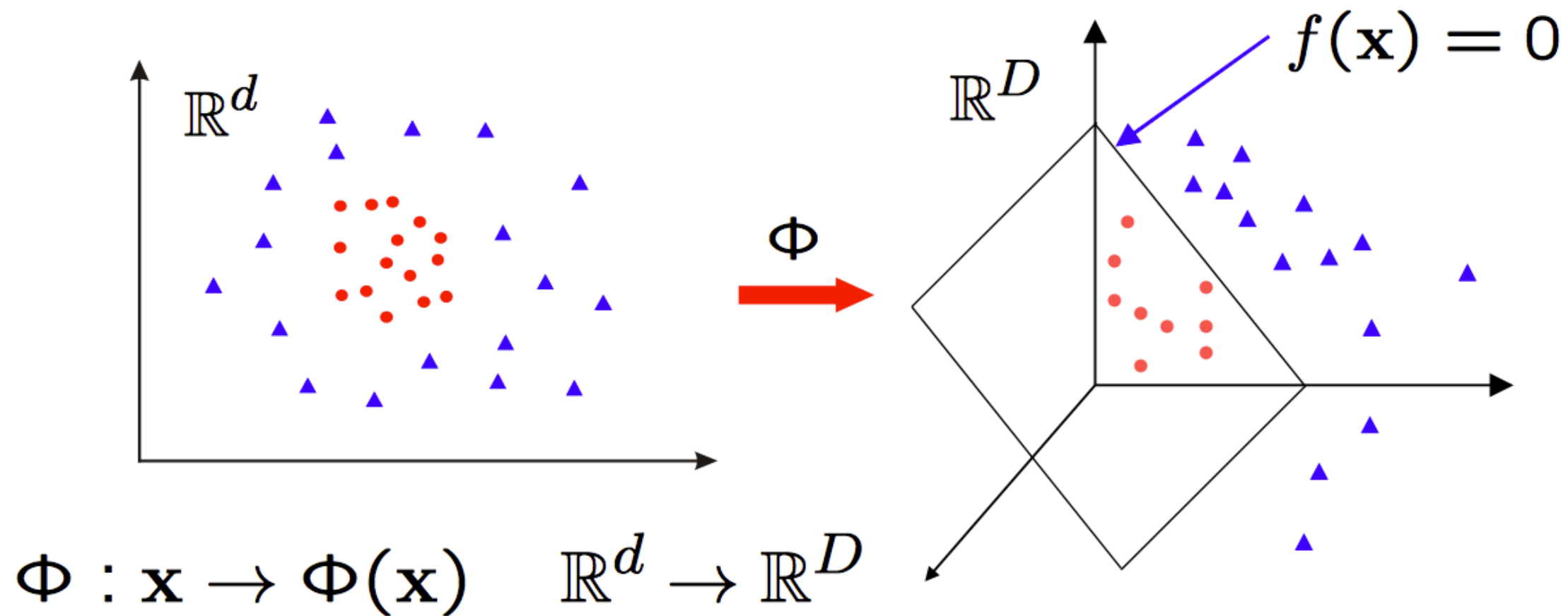
Idea 1: Map Data to Higher Dimension $\mathbb{Z}$ space

$$\Phi : \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \rightarrow \begin{pmatrix} x_1^2 \\ x_2^2 \\ \sqrt{2}x_1x_2 \end{pmatrix} \quad \mathbb{R}^2 \rightarrow \mathbb{R}^3$$



- Data **is** linearly separable in 3D
- This means that the problem can still be solved by a linear classifier

SVM in a Transformed Feature Space



Learn classifier linear in $\mathbf{w}$ for $\mathbb{R}^D$:

$$f(x) = \phi(x)\theta + \theta_0 = \overset{\text{z space}}{z}\theta + \theta_0$$

$\Phi(\mathbf{x})$ is a feature map

Kernel trick – what do we need from $\mathbb{Z}$ space

$$\begin{aligned} & X X^T_{N \times N} \\ & X_{N \times d} \\ & \downarrow \\ & z_{N \times D} \\ & z z^T_{N \times N} \end{aligned}$$

$D \rightarrow d$

lagrange multiplier

$$l(\alpha) = \sum_{i=1}^N \alpha^{(i)} - \frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N y^{(i)} y^{(j)} \alpha^{(i)} \alpha^{(j)} \underbrace{z^{(i)} z^{(j)T}}_{\text{Inner products}}$$

Constraints: $\alpha^{(i)} \geq 0$ for $i = 1, \dots, N$ and $\sum_{i=1}^N \alpha^{(i)} y^{(i)} = 0$

We already have this:

$$\begin{bmatrix} y^{(1)} y^{(1)} & \underbrace{z^{(1)} z^{(1)T}}_{1 \times D, D \times 1} & \dots & y^{(1)} y^{(N)} & z^{(1)} z^{(N)T} \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ y^{(N)} y^{(1)} & \underbrace{z^{(N)} z^{(1)T}}_{D \times 1, 1 \times D} & \dots & y^{(N)} y^{(N)} & z^{(N)} z^{(N)T} \end{bmatrix}$$

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

Same result as hard SVM:

Solve $\alpha^{(i)}$ using quadratic programming and predict a test data point z in z space $\rightarrow$

$$\text{sign}(z\theta + b) = \sum_{z_i \text{ in } SV} \alpha^{(i)} y^{(i)} z^{(i)} z + b$$

Generalized inner product

Given **two points** $x^{\{1\}}$ and $x^{\{2\}}$
$$\begin{bmatrix} y^{\{1\}}y^{\{1\}}K(x^{\{1\}}, x^{\{1\}}) & y^{\{1\}}y^{\{2\}}K(x^{\{1\}}, x^{\{2\}}) \\ y^{\{2\}}y^{\{1\}}K(x^{\{2\}}, x^{\{1\}}) & y^{\{2\}}y^{\{2\}}K(x^{\{2\}}, x^{\{2\}}) \end{bmatrix}$$

How can we calculate $z^{\{1\}}z^{\{2\}T}$

Let $K(x^{\{1\}}, x^{\{2\}}) = z^{\{1\}}z^{\{2\}T}$ **The kernel** inner product of $x^{\{1\}}$ and $x^{\{2\}}$

Example:

$x^{\{1\}} = (1, h^{\{1\}}, w^{\{1\}}) \rightarrow 2nd - order \Phi$ here $x^{\{1\}}$ and $x^{\{2\}}$ have two dimensions

$x^{\{2\}} = (1, h^{\{2\}}, w^{\{2\}}) \rightarrow 2nd - order \Phi$

$$z^{\{1\}} = \Phi(\mathbf{x}) = (1, h^{\{1\}}, w^{\{1\}}, h^{\{1\}^2}, w^{\{1\}^2}, h^{\{1\}}w^{\{1\}})$$

$$z^{\{2\}} = \Phi(\mathbf{x}') = (1, h^{\{2\}}, w^{\{2\}}, h^{\{2\}^2}, w^{\{2\}^2}, h^{\{2\}}w^{\{2\}})$$

How many dimensions z has?

$$K(x^{\{1\}}, x^{\{2\}}) = z^{\{1\}}z^{\{2\}T} = 1 + h^{\{1\}}h^{\{2\}} + w^{\{1\}}w^{\{2\}} + h^{\{1\}^2}h^{\{2\}^2} + w^{\{1\}^2}w^{\{2\}^2} + h^{\{1\}}h^{\{2\}}w^{\{1\}}w^{\{2\}}$$

We can also calculate: $K(x^{\{1\}}, x^{\{1\}})K(x^{\{2\}}, x^{\{1\}})K(x^{\{2\}}, x^{\{2\}})$

The trick

Can we compute $K(x^{\{1\}}, x^{\{2\}})$ **without** transforming $x^{\{1\}}$ and $x^{\{2\}}$?

Example:

Datapoint 1 in **x** space

$$x^{\{1\}} = (1, h^{\{1\}}, w^{\{1\}})$$

Datapoint 2 in **x** space

$$x^{\{2\}} = (1, h^{\{2\}}, w^{\{2\}})$$

Datapoint 1 and 2 have 2 dimensions in **x** space and 1 is for the biased term

Datapoint 1 in **z** space $\phi(x^{\{1\}}) = z^{\{1\}} = (1, h^{\{1\}^2}, w^{\{1\}^2}, \sqrt{2}h^{\{1\}}, \sqrt{2}w^{\{1\}}, \sqrt{2}h^{\{1\}}w^{\{1\}})$

Datapoint 2 in **z** space $\phi(x^{\{2\}}) = z^{\{2\}} = (1, h^{\{2\}^2}, w^{\{2\}^2}, \sqrt{2}h^{\{2\}}, \sqrt{2}w^{\{2\}}, \sqrt{2}h^{\{2\}}w^{\{2\}})$

Datapoint 1 and 2 have 5 dimensions in **z** space and 1 is for the biased term

We need to calculate the dot product for the kernel

$$k(x^{(1)}, x^{(2)})$$

$$k(x^{(1)}, x^{(2)}) = z^{(1)} \cdot z^{(2)T}$$

$$K(x^{(1)}, x^{(2)}) = \phi(x^{(1)}) \phi(x^{(2)})^T = z^{(1)} z^{(2)T}$$

$$z^{(1)} = (1, h^{(1)^2}, w^{(1)^2}, \sqrt{2}h^{(1)}, \sqrt{2}w^{(1)}, \sqrt{2}h^{(1)}w^{(1)})$$

$$z^{(2)} = (1, h^{(2)^2}, w^{(2)^2}, \sqrt{2}h^{(2)}, \sqrt{2}w^{(2)}, \sqrt{2}h^{(2)}w^{(2)})$$

1

$$z^{(1)} z^{(2)T} = K(x^{(1)}, x^{(2)}) =$$

$$1 + h^{(1)^2} h^{(2)^2} + w^{(1)^2} w^{(2)^2} + 2h^{(1)} h^{(2)} + 2w^{(1)} w^{(2)} + 2h^{(1)} h^{(2)} w^{(1)} w^{(2)}$$

$$K(x^{(1)}, x^{(2)}) = (1 + h^{(1)} h^{(2)} + w^{(1)} w^{(2)})^2$$

$$x^{(1)} = (1, h^{(1)}, w^{(1)})$$

$$x^{(2)} = (1, h^{(2)}, w^{(2)})$$

$$\begin{matrix} x^{(1)} & \cdot & x^{(2)T} \\ \begin{bmatrix} 1 & h^{(1)} & w^{(1)} \end{bmatrix} & & \begin{bmatrix} 1 \\ h^{(2)} \\ w^{(2)} \end{bmatrix} \end{matrix}$$

$$z^{(1)} z^{(2)T} = K(x^{(1)}, x^{(2)}) = (x^{(1)} x^{(2)T})^2 = \text{Homogeneous kernel}$$

Example: Let's say we have two datapoints

We need to build the inner product matrix to optimize SVM parameters.

$$l(\alpha) = \sum_{i=1}^N \alpha^{\{i\}} - \frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N y^{\{i\}} y^{\{j\}} \alpha^{\{i\}} \alpha^{\{j\}} z^{\{i\}} z^{\{j\}T} = \sum_{i=1}^N \alpha^{\{i\}} - \frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N y^{\{i\}} y^{\{j\}} \alpha^{\{i\}} \alpha^{\{j\}} K(x^{\{i\}}, x^{\{j\}})$$

$X = \begin{bmatrix} 1 & 1 & 2 \\ 1 & -1 & 3 \end{bmatrix}$ $Y = \begin{bmatrix} cat \\ dog \end{bmatrix}$ $K(x^{\{1\}}, x^{\{2\}}) = \begin{bmatrix} y^{\{1\}} y^{\{1\}} K(x^{\{1\}}, x^{\{1\}}) & y^{\{1\}} y^{\{2\}} K(x^{\{1\}}, x^{\{2\}}) \\ y^{\{2\}} y^{\{1\}} K(x^{\{2\}}, x^{\{1\}}) & y^{\{2\}} y^{\{2\}} K(x^{\{2\}}, x^{\{2\}}) \end{bmatrix}$

$Z = \begin{bmatrix} 1 & \sqrt{2}w & \sqrt{2}h & w^2 & h^2 & \sqrt{2}hw \\ 1 & \sqrt{2} & 2\sqrt{2} & 1 & 4 & 2\sqrt{2} \\ 1 & -\sqrt{2} & 3\sqrt{2} & 1 & 9 & -3\sqrt{2} \end{bmatrix}$ $z^{\{1\}}$ $z^{\{2\}}$

$K(x^{\{2\}}, x^{\{1\}}) = z^{\{2\}} \cdot z^{\{1\}T} = [1 \quad -\sqrt{2} \quad 3\sqrt{2} \quad 1 \quad 9 \quad -3\sqrt{2}]$

$\begin{bmatrix} 1 \\ \sqrt{2} \\ 2\sqrt{2} \\ 1 \\ 4 \\ 2\sqrt{2} \end{bmatrix}$

$K(x^{\{2\}}, x^{\{1\}}) = (x^{\{2\}} \cdot x^{\{1\}T})^{2000}$

The polynomial kernel

$\mathcal{X} = \mathcal{R}^d$ and $\Phi: \mathcal{X} \rightarrow \mathbb{Z}$ is polynomial of order Q

The “equivalent kernel” $= K(x^{\{1\}}, x^{\{2\}}) = \left(1 + x^{\{1\}} x^{\{2\}T}\right)^Q$ Inhomogeneous kernel

$$= \left(1 + x_1^{\{1\}} x_1^{\{2\}} + x_2^{\{1\}} x_2^{\{2\}} + \cdots + x_d^{\{1\}} x_d^{\{2\}}\right)^Q$$

Does it matter if Q is 2 or 1000?

What will happen if we have $d = 10$ and $Q = 100$ and we want to compute the inner product explicitly?

We need to calculate the inner product of two big huge ugly vectors

We only need $\mathbb{Z}$ space to exist

If $K(x^{\{1\}}, x^{\{2\}})$ is an inner product in some space $\mathbb{Z}$, we are doing good

Example:

$$K(x^{\{1\}}, x^{\{2\}}) = \exp\left(-\gamma \overbrace{\|x^{\{1\}} - x^{\{2\}}\|^2}\right) \quad \text{Radial basis kernel}$$

First thing first, this is a function of x and x'

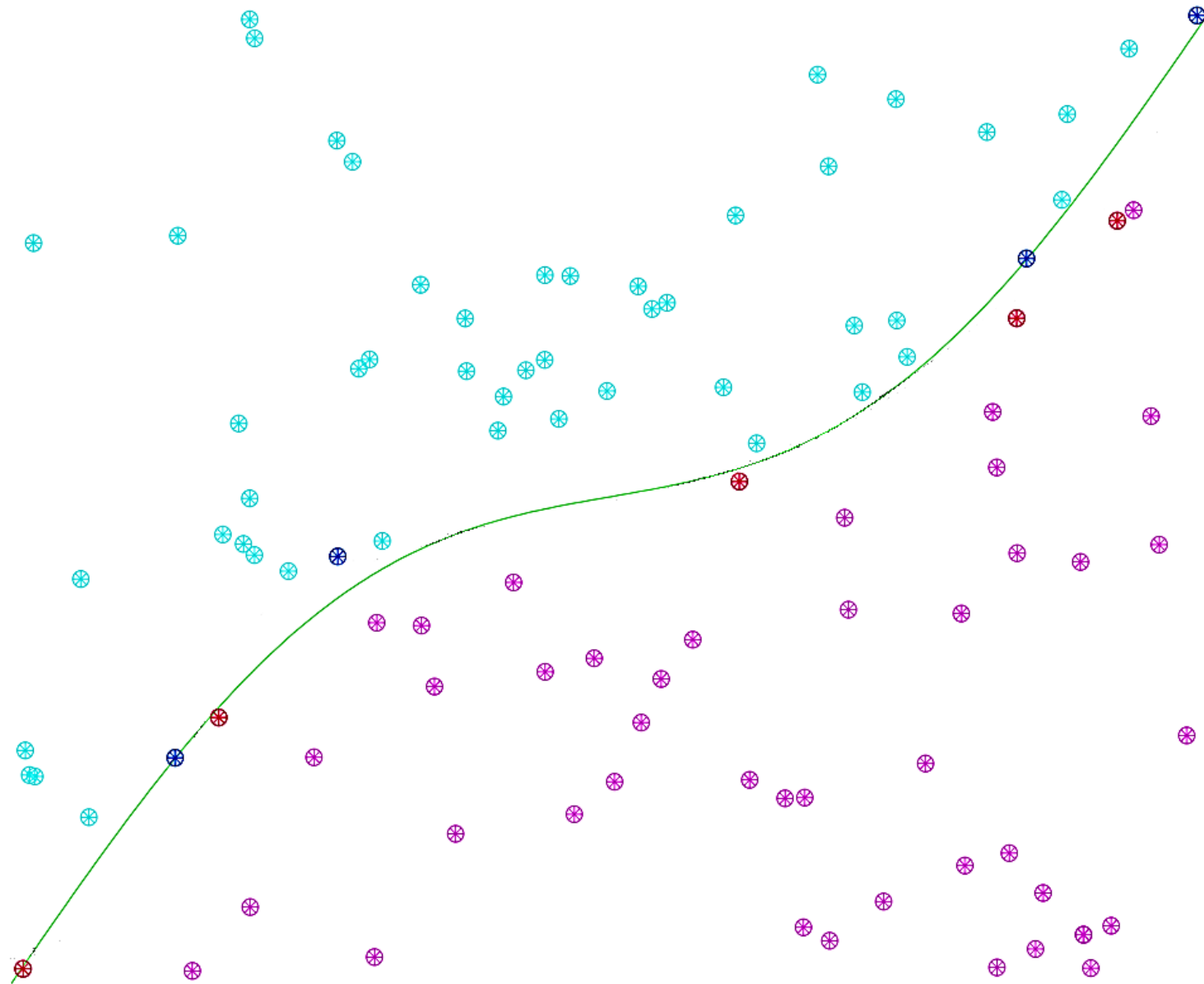
This function will take us to infinite-dimensional $\mathbb{Z} \rightarrow$ **CONGRATULATIONS**

$$\text{For } d \text{ and } \gamma = 1 \Rightarrow K(x^{\{1\}}, x^{\{2\}}) = \exp(-(x^{\{1\}} - x^{\{2\}})^2)$$

$$= \exp(-x^{\{1\}^2}) \exp(-x^{\{2\}^2}) \underbrace{\sum_{k=0}^{\infty} \frac{2^k x^{\{1\}^k} x^{\{2\}^k}}{k!}}_{\text{exp Taylor expansion for } \exp(2xx')}$$

Radial basis kernel in action

Slightly non-linearly separable case for 100 datapoints:



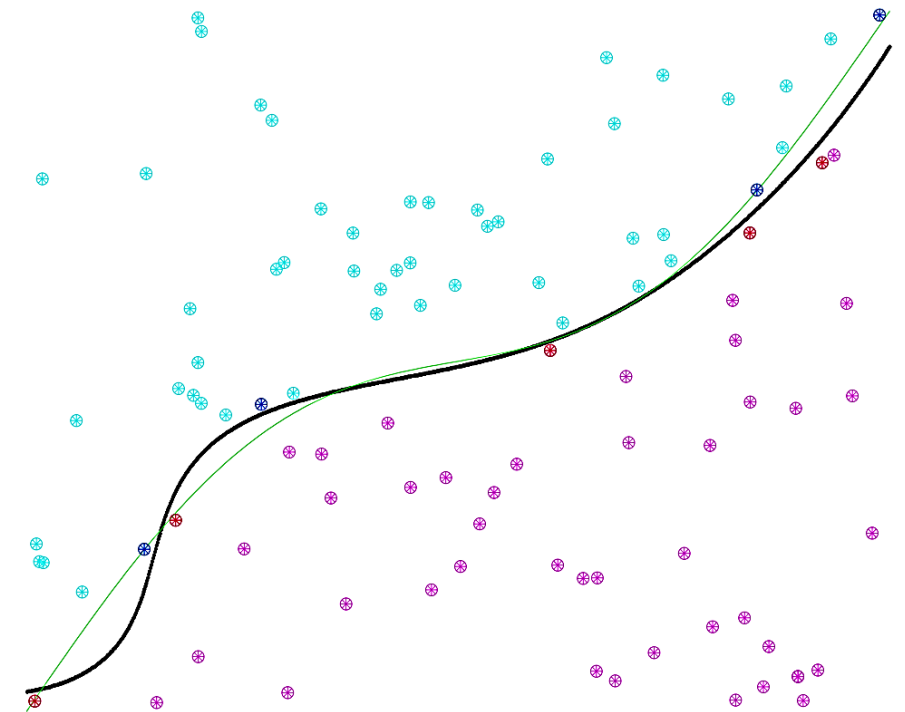
Transforming χ
into ∞ -dimensional $\mathbb{Z}$ space



Generalization

Are we killing the generalization by going to infinite-dimension? (overfitting)

I am going to answer this with a question.
In this, example how many support vectors, we have?



What will happen if we have many support vectors?

The decision boundary line (plane) will be super wiggly → overfitting alarm

$$\mathbb{E}[E_{out}] \leq \frac{\mathbb{E}[\text{Number of support vectors}]}{N - 1}$$

N is number of datapoints

Kernel formulation of SVM

Remember quadratic programming?

$$\begin{bmatrix} y^{\{1\}}y^{\{1\}}x^{\{1\}}x^{\{1\}T} & y^{\{1\}}y^{\{2\}}x^{\{1\}}x^{\{2\}T} & \dots & y^{\{1\}}y^{\{N\}}x^{\{1\}}x^{\{N\}T} \\ y^{\{2\}}y^{\{1\}}x^{\{2\}}x^{\{1\}T} & y^{\{2\}}y^{\{2\}}x^{\{2\}}x^{\{2\}T} & \dots & y^{\{2\}}y^{\{N\}}x^{\{2\}}x^{\{N\}T} \\ \dots & \dots & \dots & \dots \\ y^{\{N\}}y^{\{1\}}x^{\{N\}}x^{\{1\}T} & y^{\{N\}}y^{\{2\}}x^{\{N\}}x^{\{2\}T} & \dots & y^{\{N\}}y^{\{N\}}x^{\{N\}}x^{\{N\}T} \end{bmatrix}$$

Quadratic coefficients

In $\mathbb{Z}$ space, the only thing you need:

$$\begin{bmatrix} y^{\{1\}}y^{\{1\}}K(x^{\{1\}}, x^{\{1\}T}) & y^{\{1\}}y^{\{2\}}K(x^{\{1\}}, x^{\{2\}T}) & \dots & y^{\{1\}}y^{\{N\}}K(x^{\{1\}}, x^{\{N\}T}) \\ y^{\{2\}}y^{\{1\}}K(x^{\{2\}}, x^{\{1\}T}) & y^{\{2\}}y^{\{2\}}K(x^{\{2\}}, x^{\{2\}T}) & \dots & y^{\{2\}}y^{\{N\}}K(x^{\{2\}}, x^{\{N\}T}) \\ \dots & \dots & \dots & \dots \\ y^{\{N\}}y^{\{1\}}K(x^{\{N\}}, x^{\{1\}T}) & y^{\{N\}}y^{\{2\}}K(x^{\{N\}}, x^{\{2\}T}) & \dots & y^{\{N\}}y^{\{N\}}K(x^{\{N\}}, x^{\{N\}T}) \end{bmatrix}$$

$$X\theta + b \leftarrow x \text{ space } \hat{y}_p$$

$$\{Z\theta + b \leftarrow z \text{ space } \hat{y}_p$$

Final stage:

$$g(x) = \text{sign}(\overset{\{test\}}{z}\theta + b)$$

in terms of $K(-, -)$

$$\text{where } \theta = \sum_{z^{(i)} \text{ in SV}} \alpha^{(i)} y^{(i)} z^{(i)} \rightarrow g(x) = \text{sign}\left(\sum_{z^{(i)} \text{ in SV}} \alpha^{(i)} y^{(i)} \underbrace{z^{(i)} z^{(i)T}}_{\text{Tot}} + b\right)$$

$$g(x) = \text{sign}\left(\sum_{\alpha_i > 0} \alpha^{(i)} y^{(i)} K(x^{(i)}, \overset{\{test\}}{x}) + b\right)$$

$$\hookrightarrow K(x^{(i)}, x^{\{test\}}) = (x^{(i)} \cdot x^{\{test\}T})^{\hat{Q}}$$

$$\text{and } b: \quad b = y^{(j)} - \sum_{\alpha^{(i)} > 0, \alpha^{(j)} > 0} \alpha^{(i)} y^{(i)} K(x^{(i)}, x^{(j)})$$

for any SV $x^{(i)}$ and $x^{(j)}$

How do we know that the kernel is valid?

For a given $K(x^{\{i\}}, x^{\{j\}})$ → We can check the validity

Three approaches:

1. By construction (Polynomial one)
2. Math properties (Mercer's condition)
3. Who cares? 😊

Design your kernel

$K(x^{\{1\}}, x^{\{2\}})$ is valid *iff*

1. It is symmetric $\rightarrow K(x^{\{1\}}, x^{\{2\}}) = K(x^{\{2\}}, x^{\{1\}})$

2. The matrix $\rightarrow Q = \begin{bmatrix} y^{\{1\}}y^{\{1\}}K(x^{\{1\}}, x^{\{1\}^T}) & y^{\{1\}}y^{\{2\}}K(x^{\{1\}}, x^{\{2\}^T}) & \dots & y^{\{1\}}y^{\{N\}}K(x^{\{1\}}, x^{\{N\}^T}) \\ y^{\{2\}}y^{\{1\}}K(x^{\{2\}}, x^{\{1\}^T}) & y^{\{2\}}y^{\{2\}}K(x^{\{2\}}, x^{\{2\}^T}) & \dots & y^{\{2\}}y^{\{N\}}K(x^{\{2\}}, x^{\{N\}^T}) \\ \dots & \dots & \dots & \dots \\ y^{\{N\}}y^{\{1\}}K(x^{\{N\}}, x^{\{1\}^T}) & y^{\{N\}}y^{\{2\}}K(x^{\{N\}}, x^{\{2\}^T}) & \dots & y^{\{N\}}y^{\{N\}}K(x^{\{N\}}, x^{\{N\}^T}) \end{bmatrix}$

$$\underbrace{\begin{pmatrix} V^T \\ V \end{pmatrix}}_{1 \times N} Q \underbrace{V}_{N \times N} \underbrace{V}_{N \times 1} \geq 0$$

is positive-semi definite

For any $x^{\{1\}}, \dots, x^{\{N\}}$ (Mercer's condition)

Common Kernels

- Linear kernels

$$K(x^{\{1\}}, x^{\{2\}}) = x^{\{1\}} x^{\{2\}T}$$

- Polynomial kernels $K(x^{\{1\}}, x^{\{2\}}) = (1 + x^{\{1\}} x^{\{2\}T})^d$ for any $d > 0$

– Contains all polynomials terms up to degree d

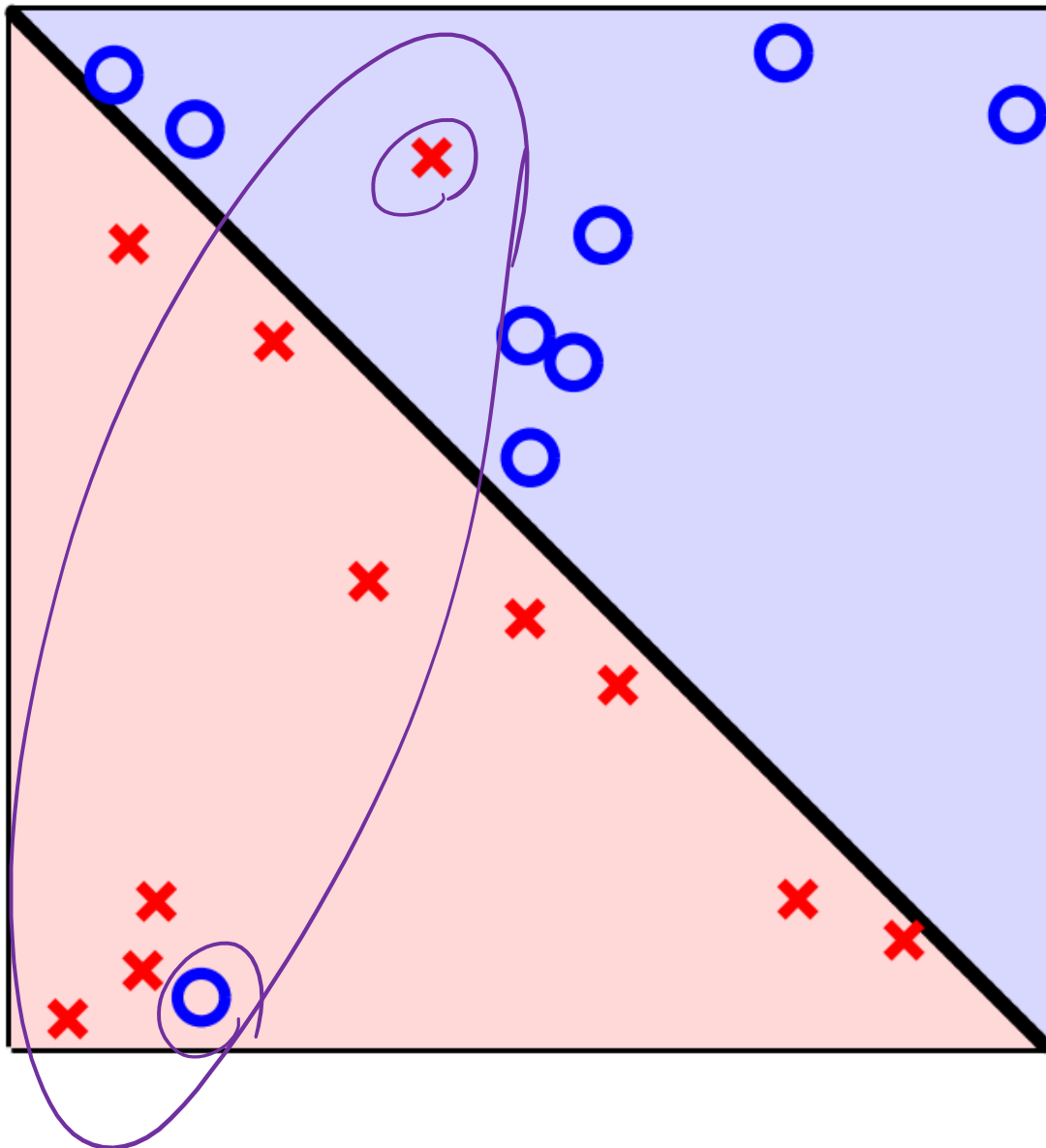
- Gaussian kernels $K(x^{\{1\}}, x^{\{2\}}) = \exp\left(-\frac{\|x^{\{1\}} - x^{\{2\}}\|^2}{2 * \sigma^2}\right)$ for $\sigma > 0$

– Infinite dimensional feature space

$$\gamma = -\frac{1}{2\sigma^2}$$

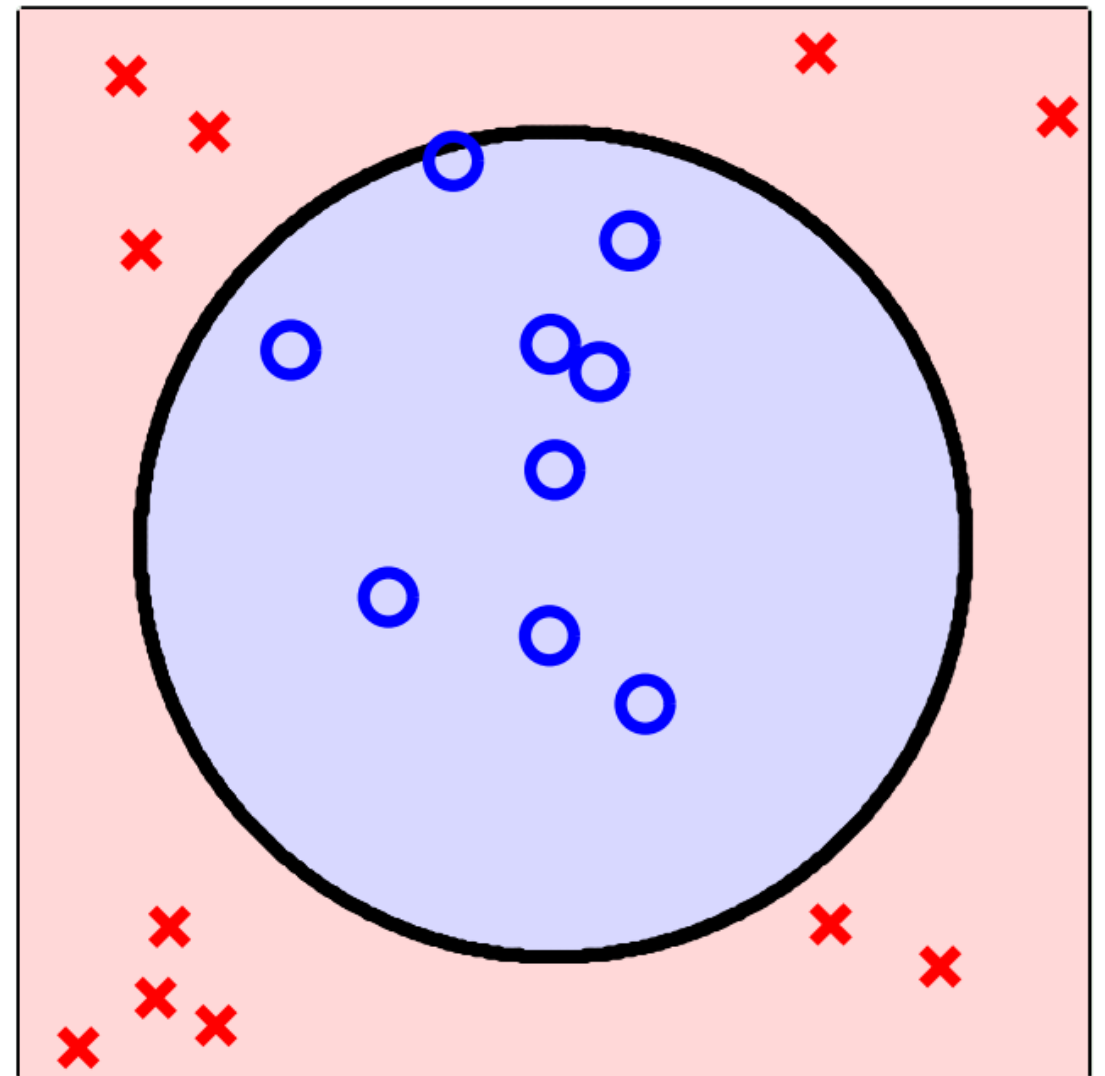
Soft SVM – Two types of non-separable

slightly:



Soft SVM will deal with this

seriously:



Kernel will deal with this



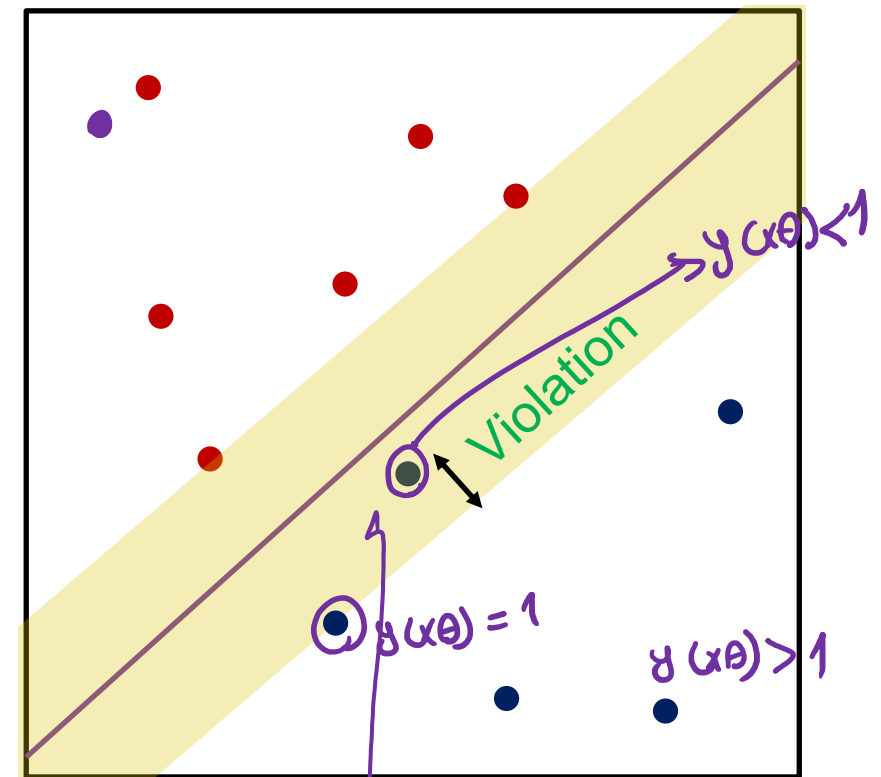
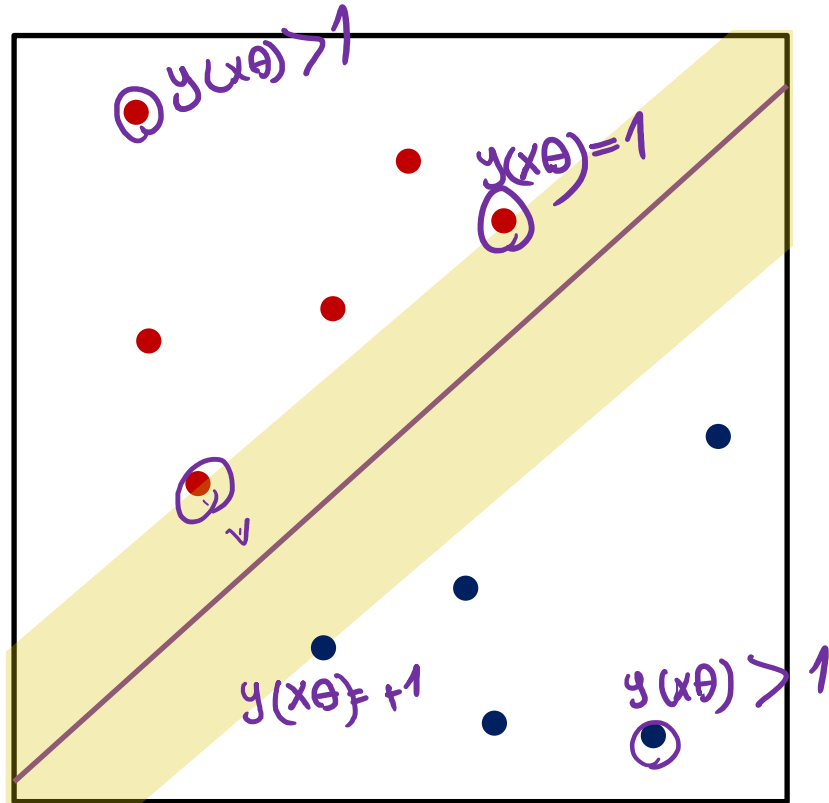
Error measure

$$\Theta \Theta^T = C$$

Non-violated case:

$$\sum_{i=1}^N (y_a^{(i)} - y_p^{(i)})^2 + \underbrace{\Theta \Theta^T}_{\text{penalty}}$$

Margin violation:



if $y^{(i)}(x^{(i)}\theta + b) > 1 \Rightarrow \text{Non SV}$

Let's introduce a slack variable: $y^{(i)}(x^{(i)}\theta + b) \geq 1 - \xi^{(i)}$

$$\xi^{(i)} \geq 0$$

Total violation = $\sum_{i=1}^N \xi^{(i)}$

$\text{Max } \frac{1}{\|\theta\|}$
 $\text{min } \frac{1}{2} \theta \theta^T$

The new optimization

Minimize $\frac{1}{2} \theta \theta^T + C \sum_{i=1}^N \xi^{(i)}$

C will define the relative importance of the first or second term

$C = \inf$ is equal to Hard SVM (will see soon)

$\alpha \rightarrow \alpha^{(1)}, \dots, \alpha^{(N)}$
 $s.t. \quad y^{(i)} (x^{(i)} \theta + b) \geq 1 - \xi^{(i)} \quad \text{for } i = 1, \dots, N$
 $and \quad \xi^{(i)} \geq 0 \quad \text{for } i = 1, \dots, N$
 $\beta^{(1)}, \dots, \beta^{(N)}$

$g(x) = y^{(i)} (x^{(i)} \theta + b) - 1 + \xi^{(i)} \quad \underline{g(x) \alpha^{(i)} = 0}$
 $h(x) = \xi^{(i)} \quad \underline{\xi^{(i)} \beta^{(i)} = 0}$

$\theta \in \mathbb{R}^d, b \in \mathbb{R}, \xi \in \mathbb{R}^N$

The Lagrange formulation

Hard svm: Minimize $\frac{1}{2} \theta \theta^T$ s.t. $y^{\{i\}} (x^{\{i\}} \theta + b) \geq 1$

$\text{Min } \mathcal{L}(\theta, b, \alpha) = \frac{1}{2} \theta \theta^T - \sum_{i=1}^N \alpha^{\{i\}} (y^{\{i\}} (x^{\{i\}} \theta + b) - 1)$

Soft svm: Minimize $\frac{1}{2} \theta \theta^T + c \sum_{i=1}^N \xi^{\{i\}}$ s.t. $y^{\{i\}} (x^{\{i\}} \theta + b) \geq 1 - \xi^{\{i\}}$ and $\xi^{\{i\}} \geq 0$

$\mathcal{L}(\theta, b, \xi, \alpha) = \frac{1}{2} \theta \theta^T + c \sum_{i=1}^N \xi^{\{i\}} - \sum_{i=1}^N \alpha^{\{i\}} (y^{\{i\}} (x^{\{i\}} \theta + b) - 1) + \sum_{i=1}^N \xi^{\{i\}} - \sum_{i=1}^N \beta^{\{i\}} \xi^{\{i\}}$

PLEASE do not scare, terms will be dropping fast

$$\mathcal{L}(\theta, b, \xi, \alpha) = \frac{1}{2} \theta \theta^T + C \sum_{i=1}^N \xi^{(i)} - \sum_{i=1}^N \alpha^{(i)} (y^{(i)} (x^{(i)} \theta + b) - 1 + \xi^{(i)}) - \sum_{i=1}^N \beta^{(i)} \xi^{(i)}$$

Handwritten notes:
 $-\sum C \xi^{(i)} + \sum \alpha^{(i)} \xi^{(i)}$ (crossed out)
 $-\sum \alpha^{(i)} \xi^{(i)} - \dots$ (crossed out)
 $\sum_{i=1}^N \overbrace{\beta^{(i)} \xi^{(i)}}^{C - \alpha^{(i)}}$ (circled)

Minimize w.r.t $\theta, b, \text{ and } \xi$ and maximize w.r.t $\alpha^{(i)} \geq 0 \text{ and } \beta^{(i)} \geq 0$

KKT condition for inequality constraints

Let's do the minimization:

$$\nabla_{\theta} \mathcal{L}(\theta, b, \xi, \alpha) = \theta - \sum_{i=1}^N \alpha^{(i)} y^{(i)} x^{(i)} = 0$$

Handwritten note: $\theta = \sum_{i=1}^N \alpha^{(i)} y^{(i)} x^{(i)}$

If we substitute $\beta^{(i)}$ up there, the whole formulation will get back to hard svm

$$\nabla_b \mathcal{L}(\theta, b, \xi, \alpha) = - \sum_{i=1}^N \alpha^{(i)} y^{(i)} = 0$$

$$\nabla_{\xi} \mathcal{L}(\theta, b, \xi, \alpha) = C - \alpha^{(i)} - \beta^{(i)} = 0$$

$$\Rightarrow \beta^{(i)} = C - \alpha^{(i)}$$

We should say thank you $\beta^{(i)}$ for the great service

$$\beta^{(i)} \geq 0 \quad \beta^{(i)} = 0$$

$$\beta^{(i)} = C - \alpha^{(i)}$$

The solution

$$\beta^{(i)} \geq 0 \Rightarrow C - \alpha^{(i)} \geq 0 \Rightarrow \alpha^{(i)} \leq C$$

$$\alpha^{(i)} g(x) = 0$$

$$\alpha^{(i)} \geq 0, \quad 0 \leq \alpha^{(i)} \leq C$$

$$\beta^{(i)} \geq 0 \quad \Rightarrow \quad C - \alpha^{(i)} \geq 0 \rightarrow 0 \leq \alpha^{(i)} \leq C \quad \text{for } i = 1, \dots, N$$

$$\text{Maximize } \mathcal{L}(\alpha) = \sum_{i=1}^N \alpha^{(i)} - \frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N y^{(i)} y^{(j)} \alpha^{(i)} \alpha^{(j)} x^{(i)} x^{(j)T} \quad \text{w.r.t } \alpha$$

$$\text{s.t.} \quad 0 \leq \alpha^{(i)} \leq C \quad \text{for } i = 1, \dots, N \quad \text{and} \quad \sum_{i=1}^N \alpha^{(i)} y^{(i)} = 0$$

$$\Rightarrow \theta = \sum_{i=1}^N \alpha^{(i)} y^{(i)} x^{(i)} \quad \text{will minimize} \quad \frac{1}{2} \theta \theta^T + C \sum_{i=1}^N \xi^{(i)}$$

Type of support vectors

$$\xi = 0 \quad \beta = 0 \Rightarrow \beta > 0 \Rightarrow C - \alpha > 0 \Rightarrow \alpha < C$$

$$0 < \alpha < C$$

We call the three points as **margin** support vectors

$$0 < \alpha^{i} < C$$

$$\beta^{i} = C - \alpha^{i}$$

$$y^{i}(x^{i}\theta + b) = 1 \Rightarrow \beta^{i} > 0 \Rightarrow \xi^{i} = 0 \text{ (KKT condition)}$$

$$\beta = 0 \quad \beta = 0 = C - \alpha = 0 \Rightarrow \alpha = C$$

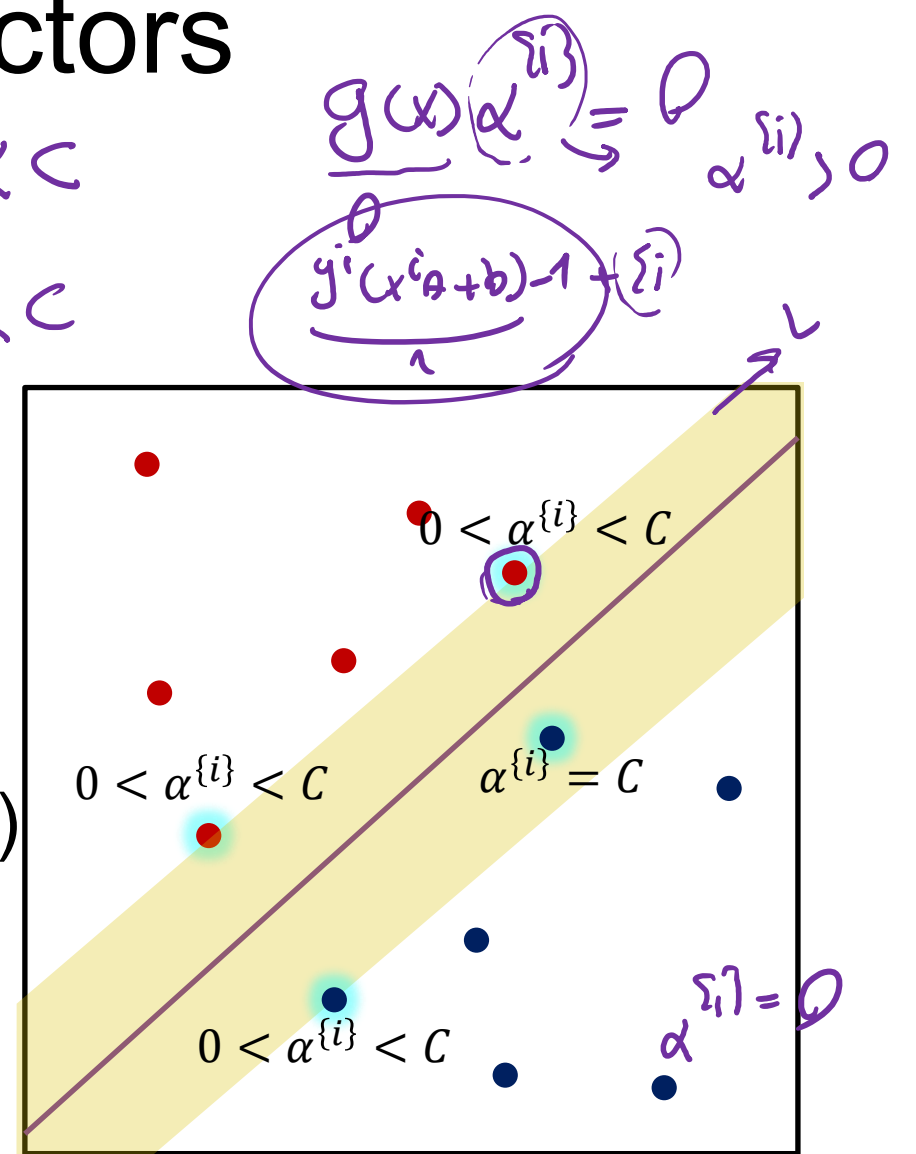
non-margin support vectors ($\alpha^{i} = C$)

$$\beta^{i} = 0 \Rightarrow \xi^{i} > 0 \text{ (KKT condition)}$$

$$y^{i}(x^{i}\theta + b) > 1 - \xi^{i} \text{ if } \xi^{i} > 0$$

$$y^{i}(x^{i}\theta + b) < 1$$

Any violating points become support vectors



$$\alpha^{i} = 0 \Rightarrow y^{i}(x^{i}\theta + b) > 1$$

Non SV

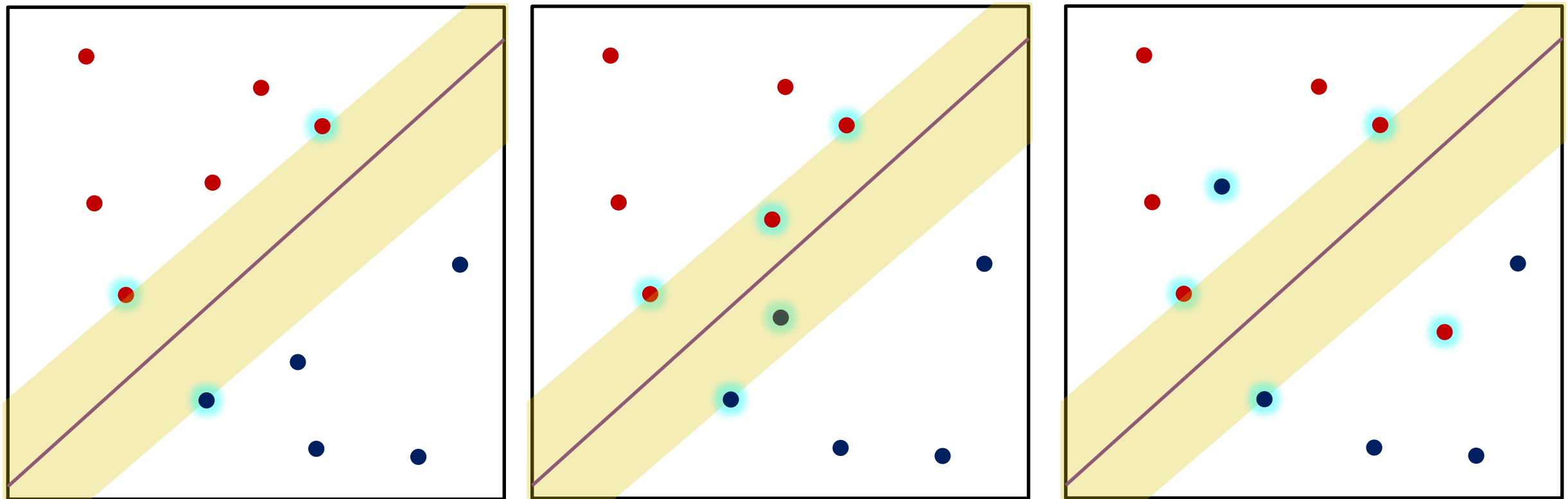
$$\alpha^{i} = C \Rightarrow y^{i}(x^{i}\theta + b) < 1$$

SV on the wrong side

$$0 < \alpha^{i} < C \Rightarrow y^{i}(x^{i}\theta + b) = 1$$

SV on the margin

How to choose C?



violating points become
support vectors

How to define the hyper-parameter C:

Cross Validation

Primal and Dual Forms of SVM

Primal version of classifier:

$$f(x_{test}) = x_{test}\theta + \theta_0$$

Dual version of classifier:

$$f(x_{test}) = \sum_{x^{(i)} \text{ in } SV} \alpha^{(i)} y^{(i)} x^{(i)} x_{test}^T + b$$

$\langle x^{(i)}, x_{test} \rangle$

Soft kernel SVM

Kernel SVM: Summary

- Classifiers can be learnt for high dimensional features spaces, without actually having to map the points into the high dimensional space
- Data may be linearly separable in the high dimensional space, but not linearly separable in the original feature space
- Kernels can be used for an SVM because of the scalar product in the dual form, but can also be used elsewhere – they are not tied to the SVM formalism