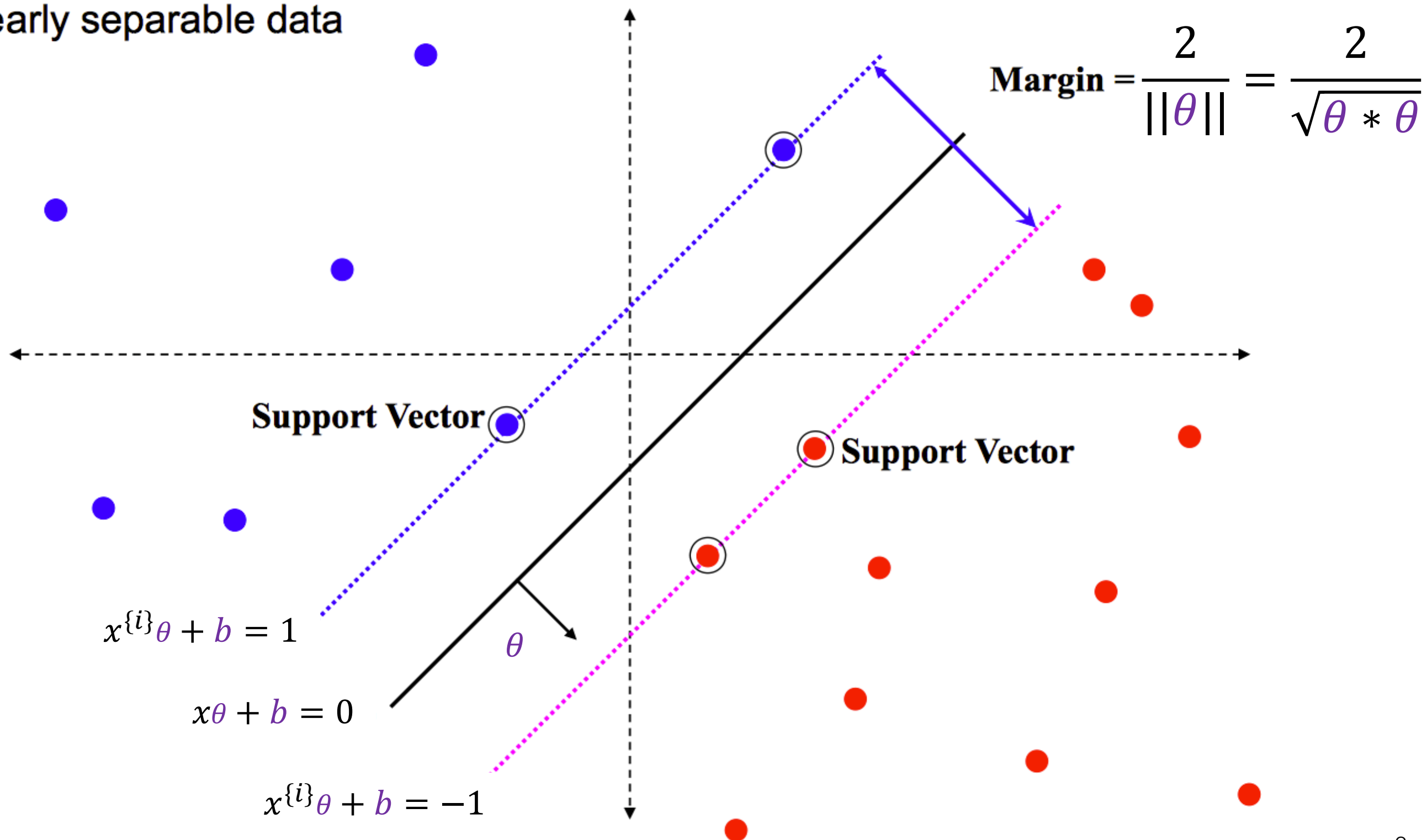


Kernel SVM

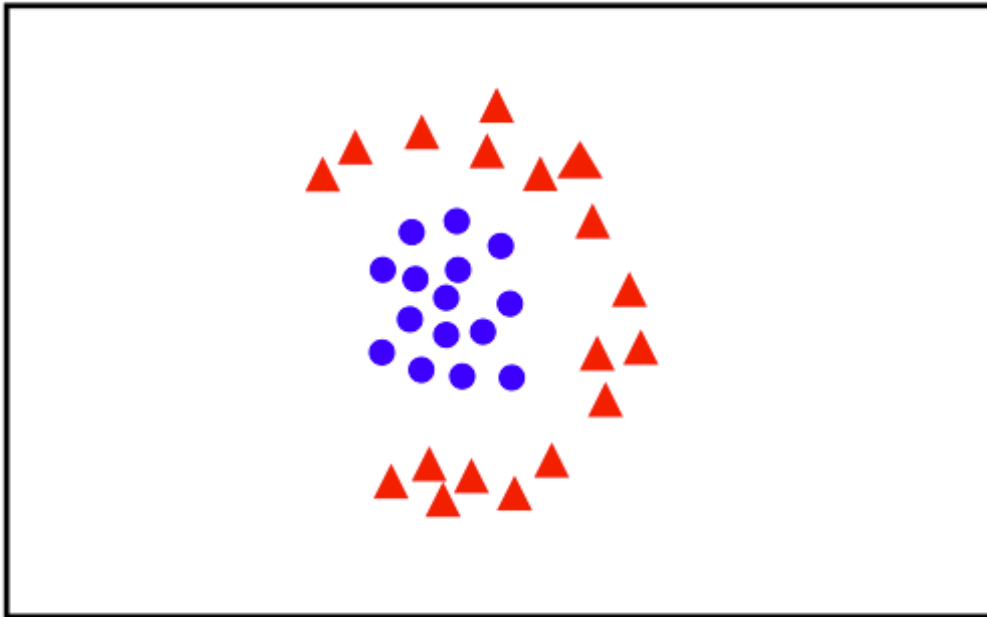
Mahdi Roozbahani
Georgia Tech

Geometric Interpretation

linearly separable data



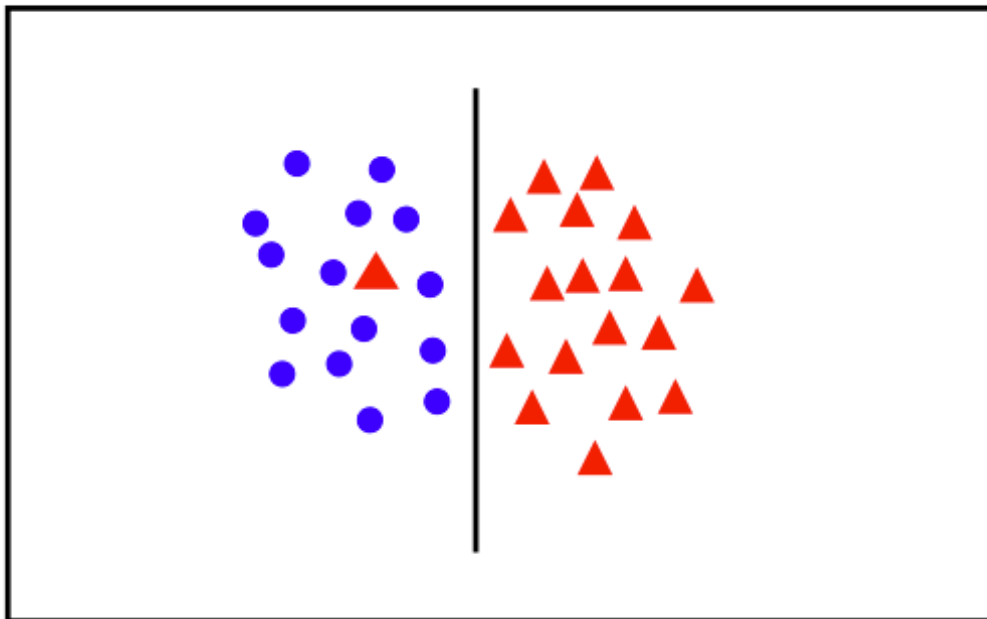
Handling Data that are Not Linearly Separable



- linear classifier not appropriate

??

Kernel trick

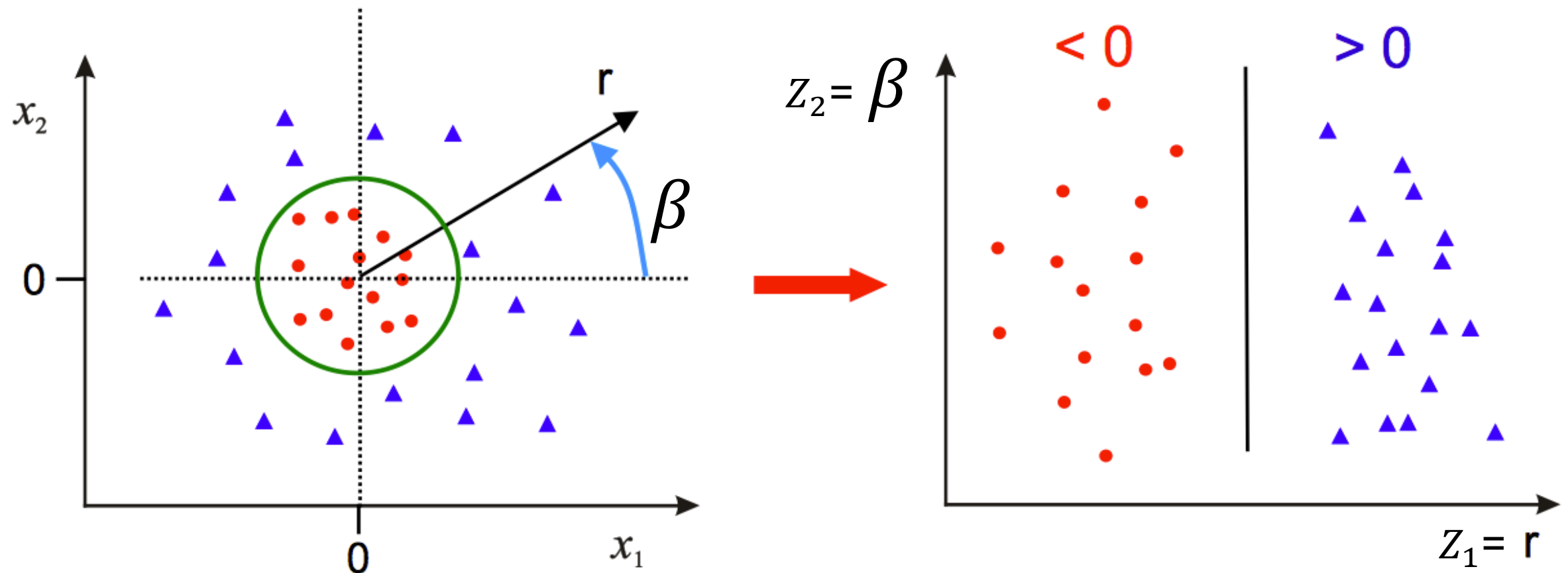


- introduce slack variables

Soft Margin SVM

(allowing ourselves to make errors)

Idea 1: Use Polar Coordinates to go to z space

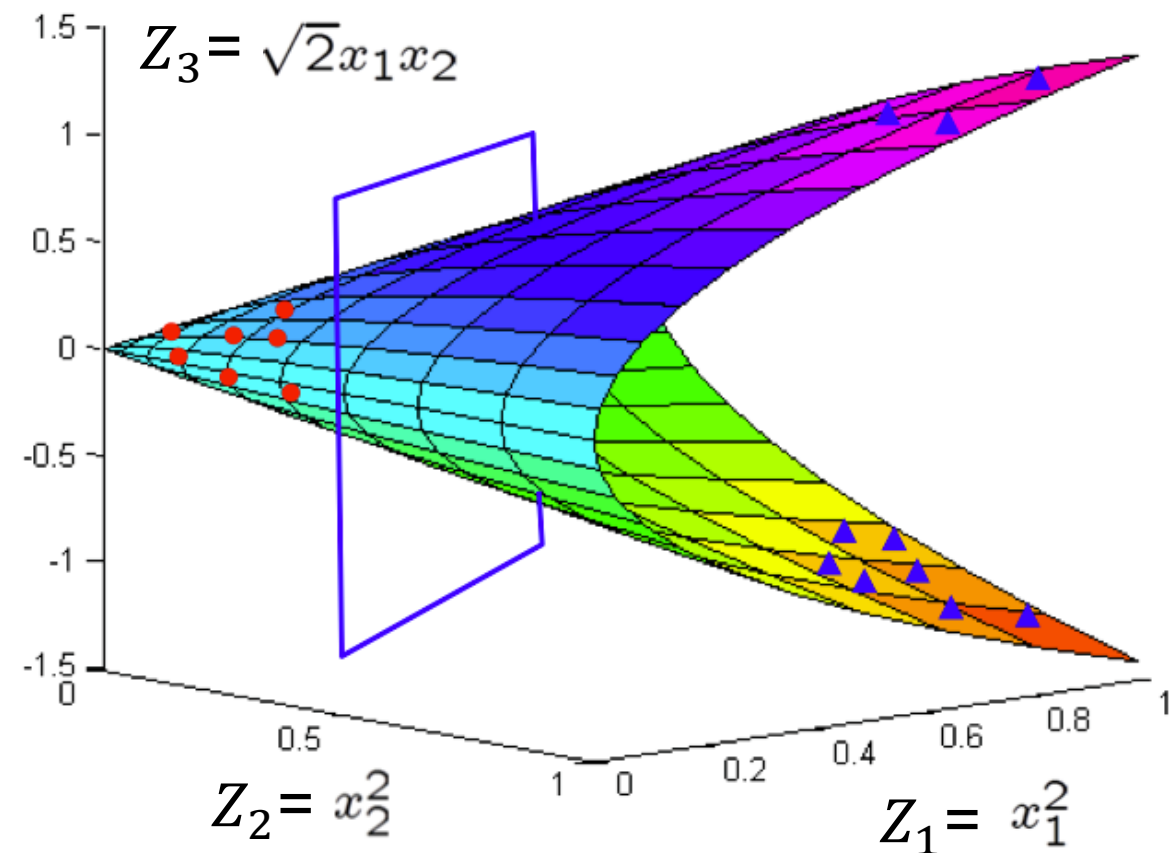
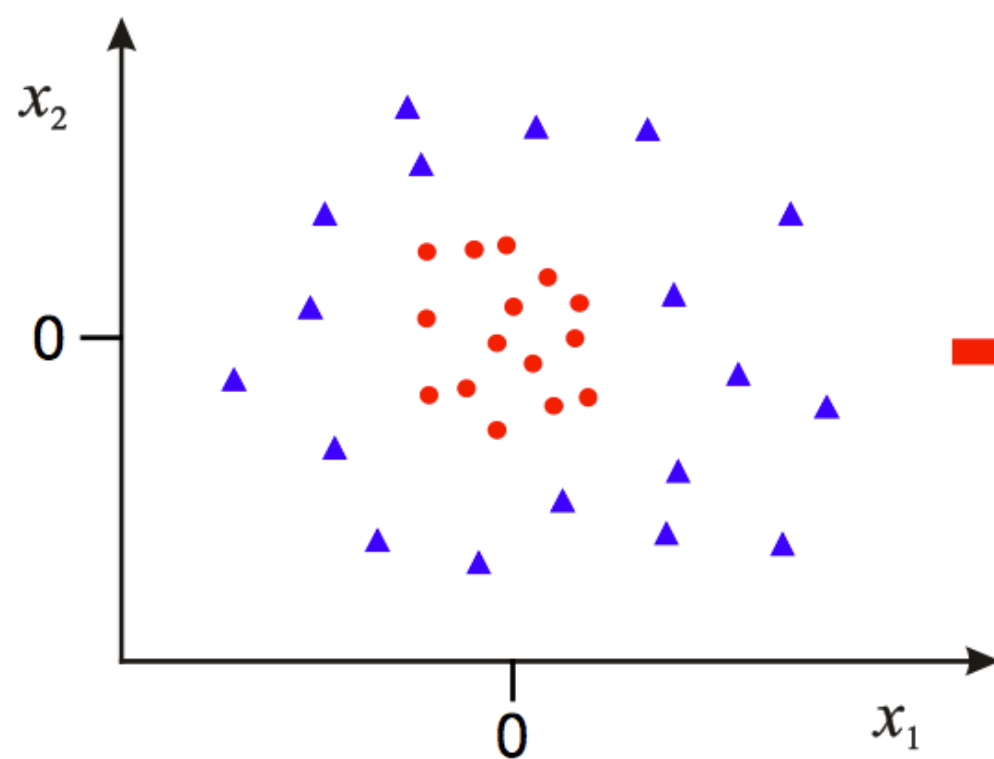


- Data **is** linearly separable in polar coordinates
- Acts non-linearly in original space

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

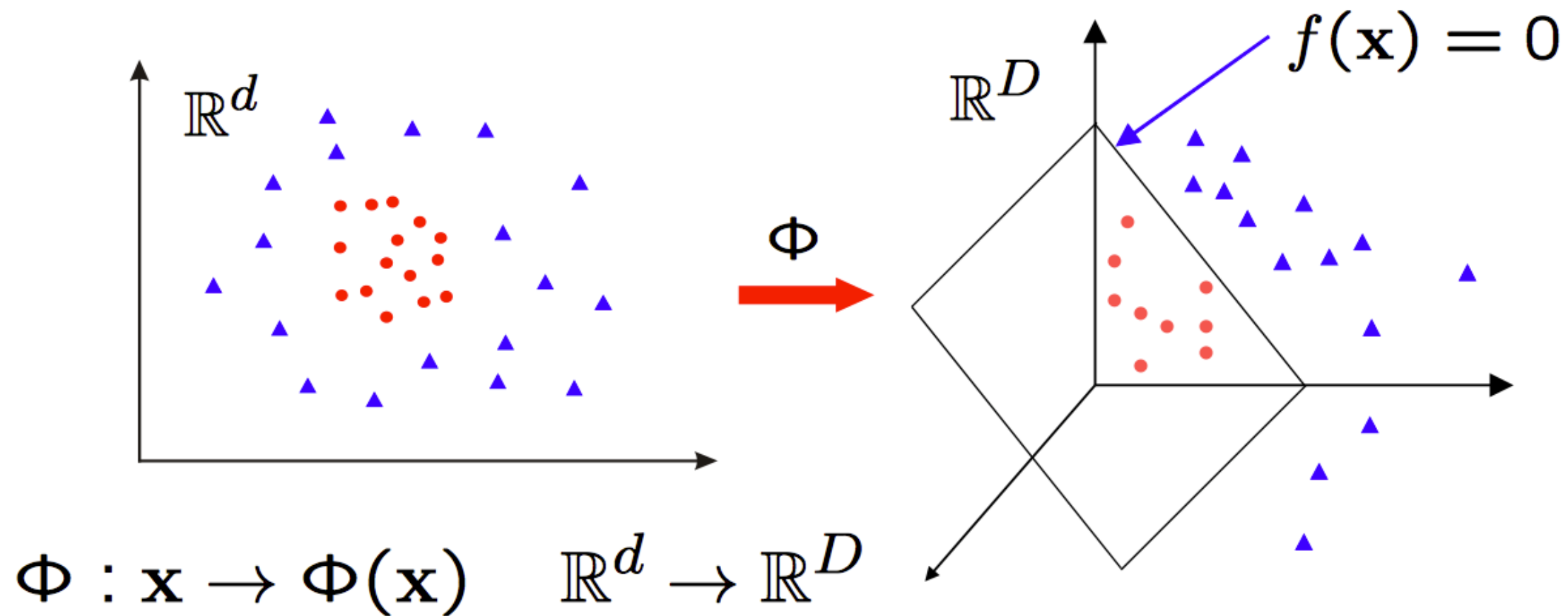
Idea 1: Map Data to Higher Dimension $\mathbb{Z}$ space

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



- 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 $\mathbb{R}^D$:

$$f(x) = \phi(x)\theta + \theta_0 = z\theta + \theta_0$$

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

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

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

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

We already have this:
$$\begin{bmatrix} y^{\{1\}} y^{\{1\}} z^{\{1\}} z^{\{1\}T} & \dots & y^{\{1\}} y^{\{N\}} z^{\{1\}} z^{\{N\}T} \\ \vdots & \ddots & \vdots \\ y^{\{N\}} y^{\{1\}} z^{\{N\}} z^{\{1\}T} & \dots & y^{\{N\}} y^{\{N\}} z^{\{N\}} z^{\{N\}T} \end{bmatrix}$$

Same result as hard SVM:

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

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

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

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

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

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

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

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

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

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

Example: Let's say we have two datapoints

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

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

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

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:

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

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

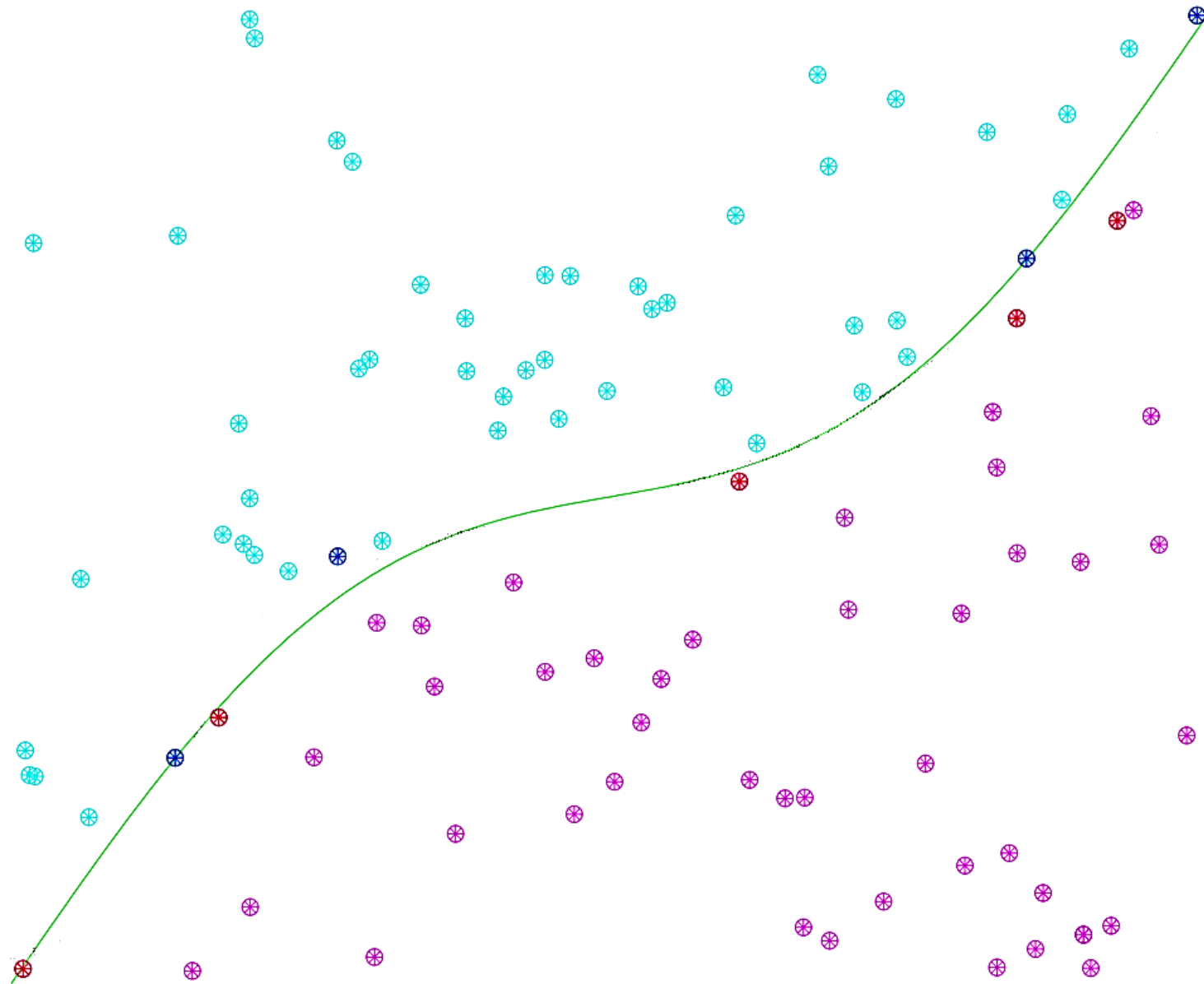
This function will take us to infinite-dimensional $\mathbb{Z} \rightarrow \text{CONGRATULATIONS}$

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

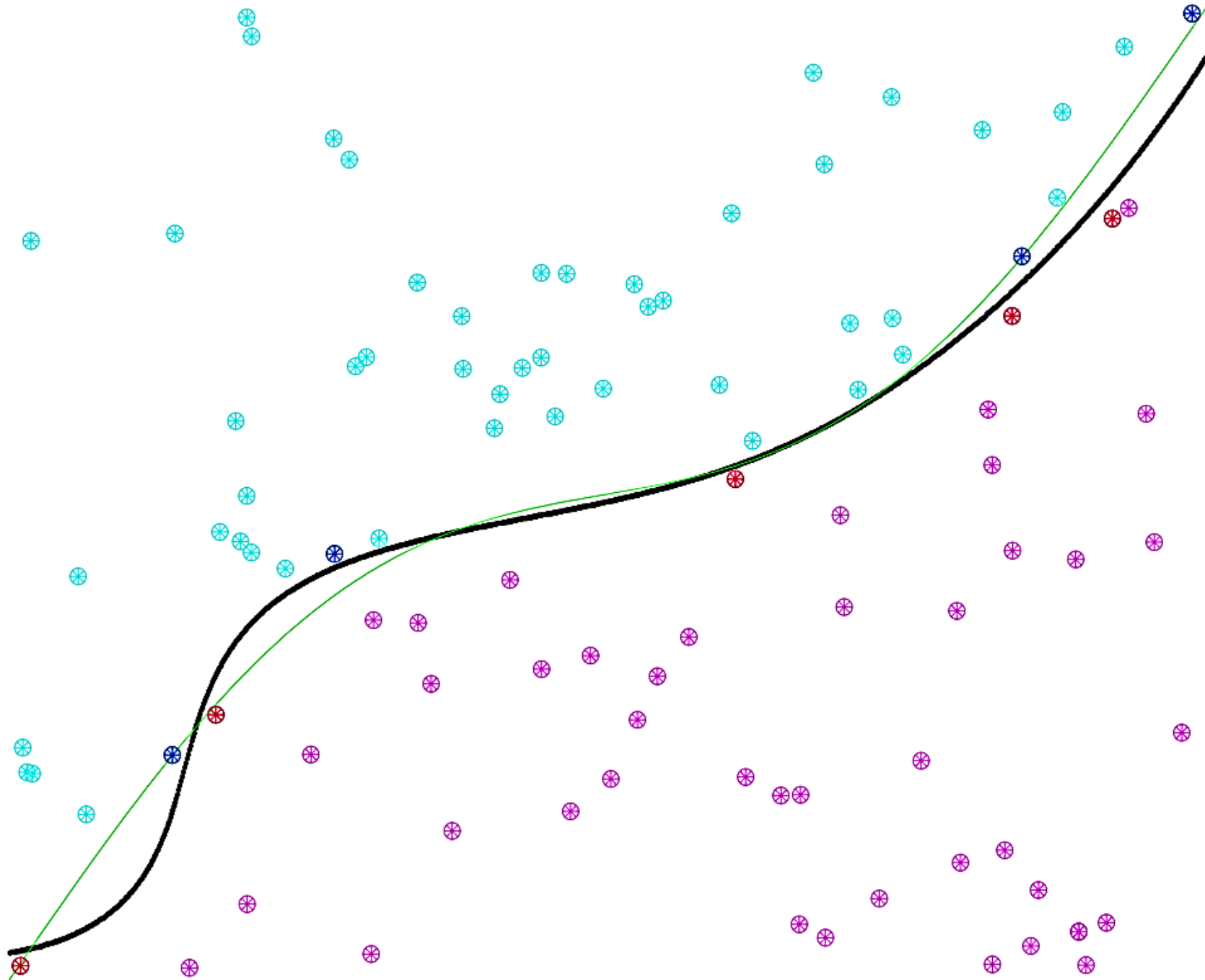
$$= \exp(-x^{\{1\}^2}) \exp(-x^{\{2\}^2}) \underbrace{\sum_{k=0}^{\infty} \frac{2^k x^{\{1\}^k} x^{\{2\}^k}}{k!}}_{\text{exp Taylor expansion for } \exp(2xx')}$$

Radial basis kernel in action

Slightly non-linearly separable case for 100 datapoints:



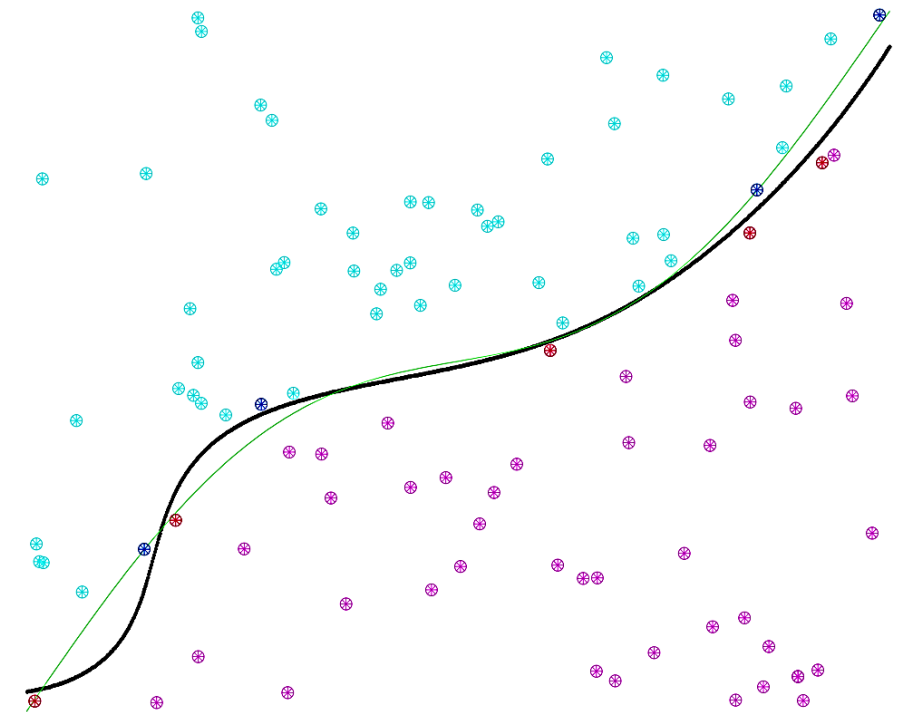
Transforming χ
into ∞ -dimensional $\mathbb{Z}$ space



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:

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

Final stage:

$$g(x) = \text{sign}(\mathbf{z}\theta + b) \quad \text{in terms of } K(-, -)$$

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

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

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

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

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 $\Rightarrow$

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

is positive-semi definite

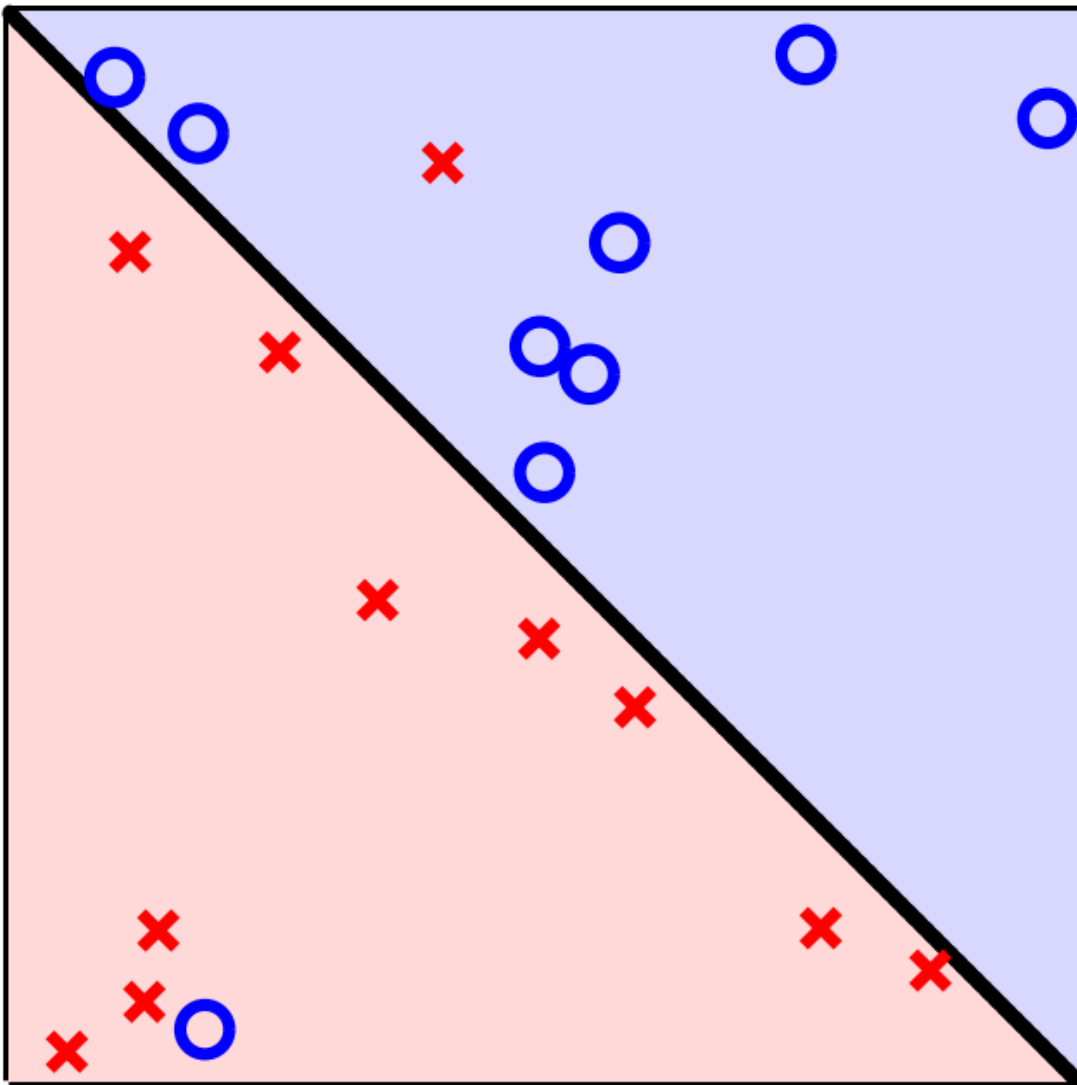
For any $x^{\{1\}}, \dots, x^{\{N\}}$ (Mercer's condition)

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

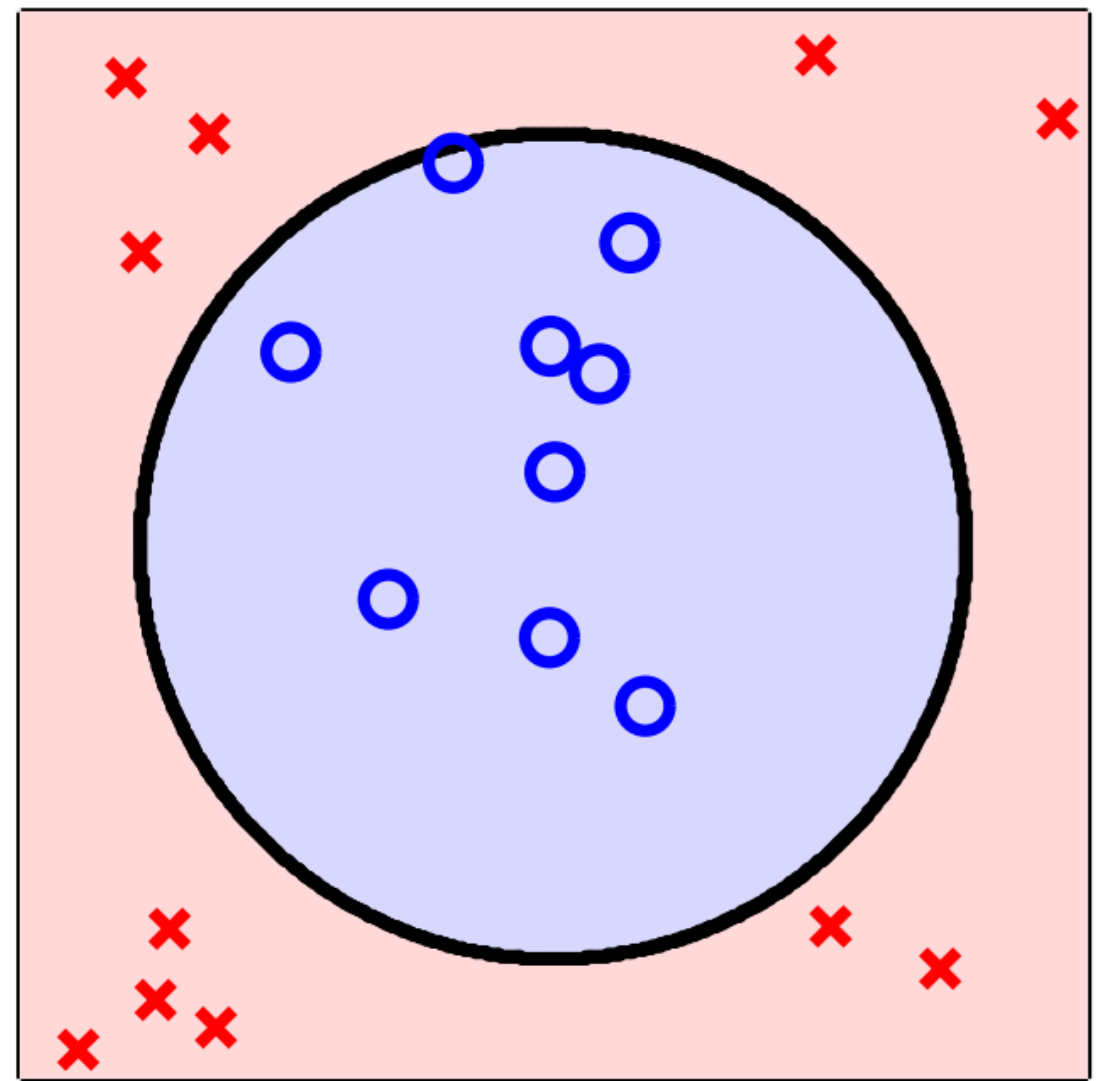
Soft SVM – Two types of non-separable

slightly:



Soft SVM will deal with this

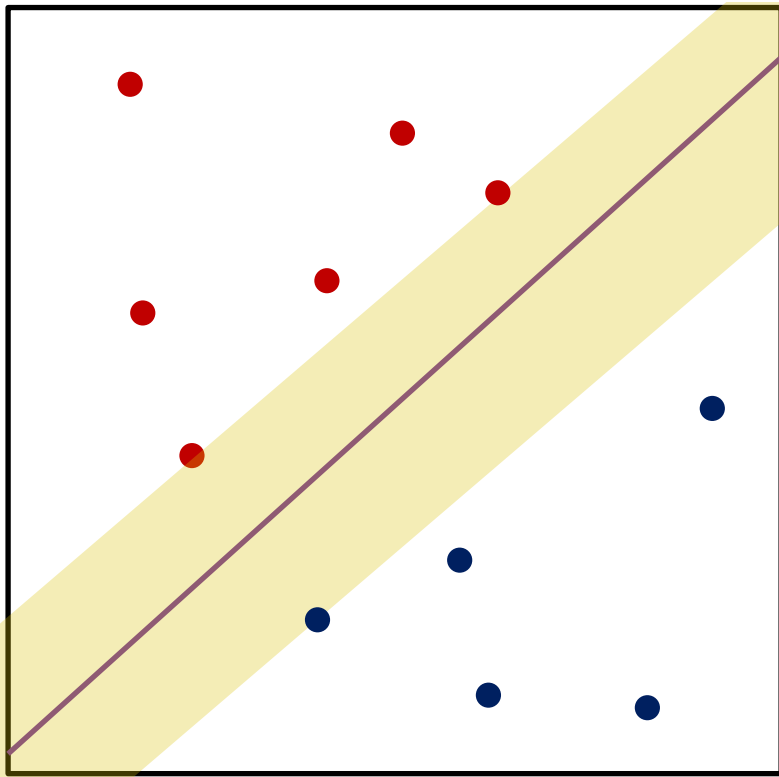
seriously:



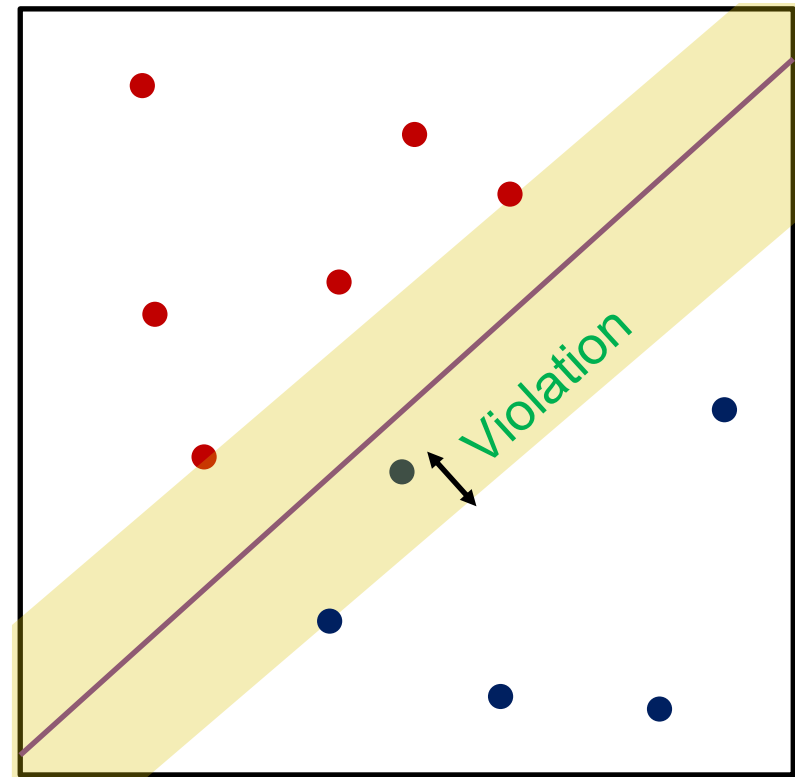
Kernel will deal with this

Error measure

Non-violated case:



Margin violation:



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

Let's introduce a slack variable: $y^{\{i\}}(x^{\{i\}}\theta + b) \geq 1 - \xi^{\{i\}}$ $\xi^{\{i\}} \geq 0$

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

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)

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

$$\text{and } \xi^{\{i\}} \geq 0 \quad \text{for } i = 1, \dots, N$$

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

The Lagrange formulation

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

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

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

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

PLEASE do not scare, terms will be dropping fast

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

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:

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

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

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

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

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

The solution

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

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

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

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

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

Type of support vectors

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

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

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

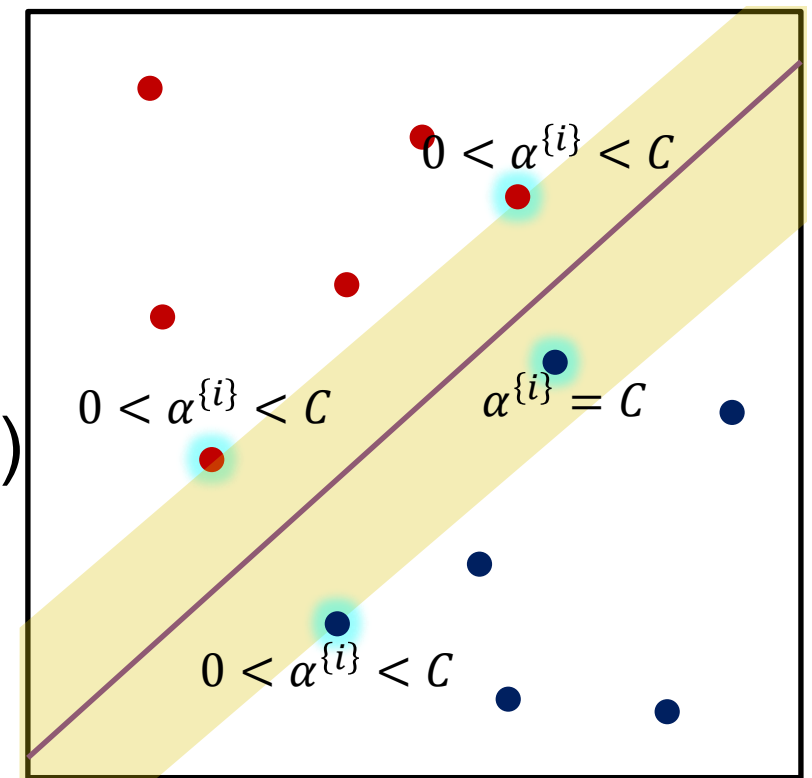
$$y^{\{i\}}(x^{\{i\}}\theta + b) = 1 \Rightarrow \beta^{\{i\}} > 0 \Rightarrow \xi^{\{i\}} = 0 \text{ (KKT condition)}$$

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

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

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

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



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

Non SV

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

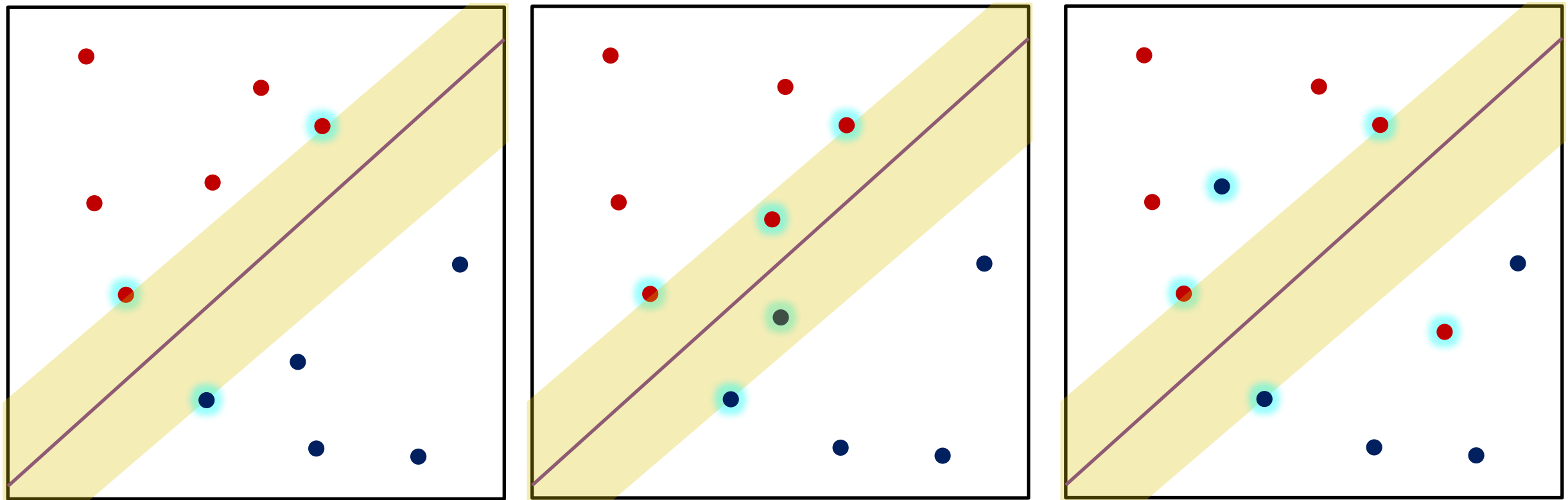
SV on the wrong side

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

SV on the margin

Any violating points become support vectors

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:

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

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