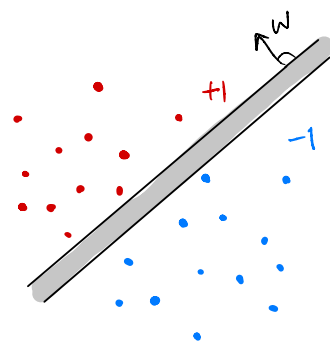


## LECTURE 20

- Last class: Binary Classification Problem

test results  $\downarrow$  diagnosis  
Training data:  $(x_i, y_i)$ ,  $i=1, \dots, n$ .  
 $\uparrow \quad \uparrow$   
 $\mathbb{R}^d \quad \{-1, 1\}$



Learning: find  $w \in \mathbb{R}^d$ :  $\langle w, x_i \rangle \begin{cases} > 1 & \text{if } y_i = 1 \\ < -1 & \text{if } y_i = -1 \end{cases}$

$$\Leftrightarrow \langle w, x_i \rangle y_i > 1 \quad \forall i$$

Prediction: diagnosis for a new  $x$ :

$$f(x) = \text{sign} \langle w, x \rangle$$

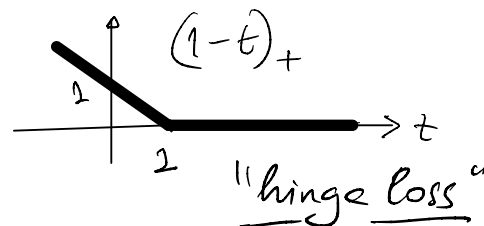
- Typical data is not perfectly separable:  $\exists$  bad pts, outliers.

To allow them, we

- Penalize: pay penalty  $1 - \langle w, x_i \rangle y_i$  for each bad pt:

$$\Rightarrow \text{loss} \quad l(w) := \sum_{i=1}^n (1 - \langle w, x_i \rangle y_i)_+ \quad \text{where}$$

Minimize  $l(w)$  over  $w \in \mathbb{R}^n$  (convex program)



- Furthermore, regularize (for robustness):

$$l(w) = \sum_{i=1}^n (1 - \langle w, x_i \rangle y_i)_+ + \lambda \|w\|_2^2$$

"Soft-margin SVM"

- And, since data may not be linearly separated, kernelize

## KERNEL METHOD (Lec. 19)

- Rewrite the algorithm in terms of inner products :  $W = \sum \alpha_j x_j \Rightarrow$

loss 
$$l(w) = \sum_{i=1}^n \left( 1 - \sum_{j=1}^n \alpha_j \langle x_i, x_j \rangle y_i \right)_+ + \lambda \sum_{j,j'=1}^n \alpha_j \alpha_{j'} \langle x_j, x_{j'} \rangle,$$

which we minimize in  $(\alpha_j)$ .

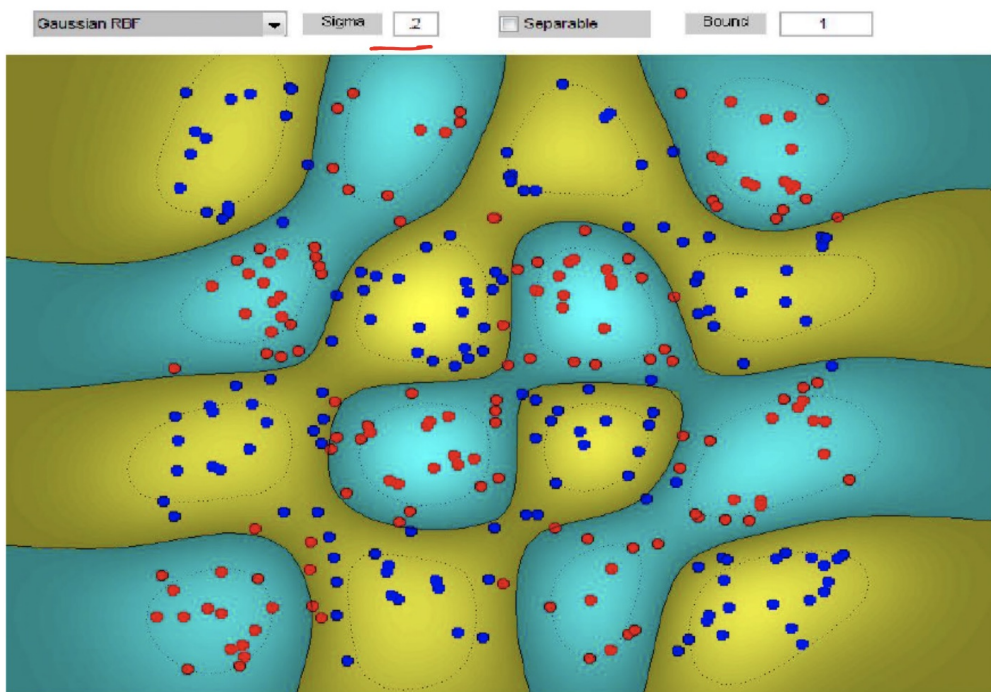
prediction:  $f(x) = \text{sign} \left( \sum \alpha_j \langle x, x_j \rangle \right)$

- Transform data using a nonlinear feature map  $\phi: \mathbb{R}^d \rightarrow H$  (implicit!)  
 $\Rightarrow \langle x, y \rangle \mapsto \langle \phi(x), \phi(y) \rangle =: K(x, y)$  "kernel"  $\uparrow$  Hilbert space
- Choose a nice kernel, e.g.  $K(x, y) = \exp\left(-\frac{\|x-y\|_2^2}{2\sigma^2}\right)$  (RBF)  
and replace all  $\langle \cdot, \cdot \rangle$  by  $K(\cdot, \cdot)$ .

"Kernel SVM"

- Warning: for practical implementation, see Wikipedia (kernelize the dual program)

EXAMPLE:



## WHAT FUNCTIONS $K(x,y)$ ARE ALLOWED?

- Can we express  $\forall$  function  $K$  as  $K(x,y) = \langle \phi(x), \phi(y) \rangle$ ?
- No: the matrix  $[K(x_i, x_j)]_{i,j=1}^n = [\langle \phi(x_i), \phi(x_j) \rangle]_{i,j=1}^n$  is a Gram matrix  $\Rightarrow$  must be PSD.

DEF A kernel is a symmetric function  $K: \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}$  such that

$$[K(x_i, x_j)]_{i,j=1}^n \succeq 0 \quad \forall x_1, \dots, x_n \in \mathbb{R}^d.$$

- This is a sufficient condition.

THM (Mercer's Condition)  $\forall$  continuous kernel  $K$

$\exists$  map  $\phi: \mathbb{R}^d \rightarrow H$  ( $H$ =Hilbert space) such that

$K(x,y) = \langle \phi(x), \phi(y) \rangle_H \quad \forall x,y \in \mathbb{R}^d$

*"feature map"* *"feature space"*

This is an  $\infty$ -dimensional version of the HW problem:

$\forall$  PSD matrix is a Gram matrix ( $x \mapsto i, y \mapsto j$ )

- Mercer's condition (PSD) is difficult to check.

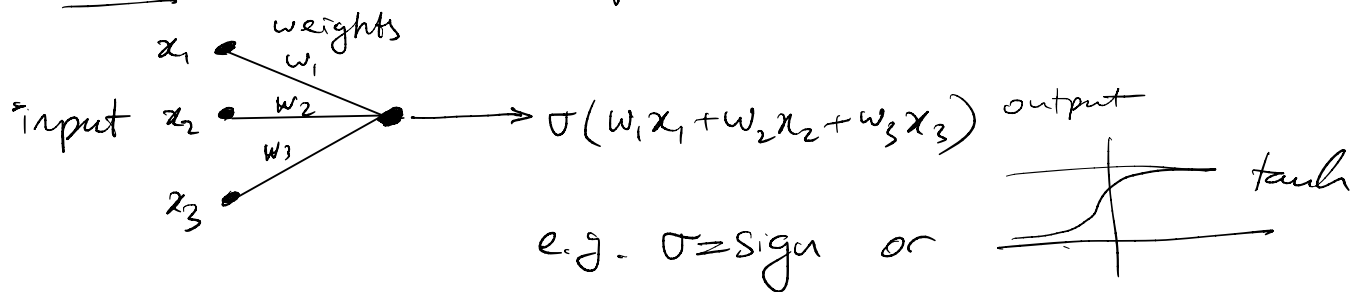
However, using it, we can build new kernels from old ones using simple rules (HW7), such as addition, exponentiation, etc.

- Examples: (a) RBF  $K(x,y) = \exp\left(-\frac{\|x-y\|_2^2}{2}\right)$ ,  
(b) polynomial  $K(x,y) = (1 + \langle x,y \rangle)^p, \dots$

Kernels arise in:

# NEURAL NETWORKS

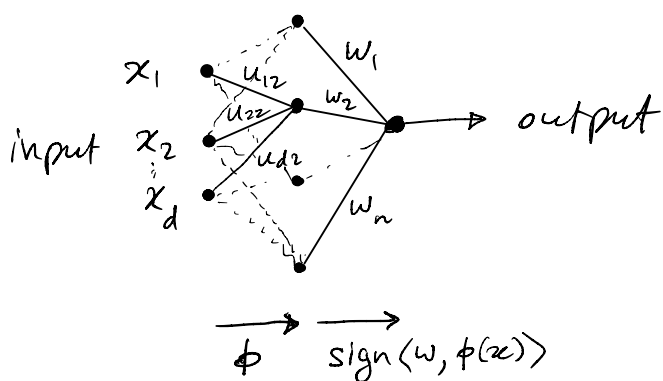
Def A neuron is a composition of linear & nonlinear function



Ex: SVM is a neuron:  $\text{sign}(w, x) = \text{sign}(w_1 x_1 + \dots + w_d x_d)$   
We train weights  $w_j$

Def A neural network is a superposition of neurons.

We train weights  $(u_{ij}), (w_j)$ .



Connection to kernels:

The first layer computes a feature map  $\phi: \mathbb{R}^d \rightarrow \mathbb{R}^n$ :

$$\phi(x) = \begin{bmatrix} \text{sgn}(u_{12}, x) \\ \vdots \\ \text{sgn}(u_{n2}, x) \end{bmatrix}$$

• Kernel:

$$K(x, y) = \langle \phi(x), \phi(y) \rangle = \frac{1}{n} \sum_{i=1}^n \text{sgn}(u_i, x) \langle u_i, y \rangle \Leftrightarrow$$

Let's initialize all weights  $u_{ij} \sim N(0, 1)$  iid

$$\Leftrightarrow \mathbb{E} \text{sgn}(u, x) \text{sgn}(u, y) \text{ where } u \sim N(0, I_n)$$

$$= \boxed{\frac{2}{\pi} \arcsin \langle x, y \rangle}$$

Hence: a neural network  $\approx$  kernel SVM.

• But: after initialization, a neural network is training weights  $u_{ij}$   
 $\Rightarrow$  trains kernel  $K$ .

• Dynamics of the prediction function is described by the neural tangent kernel (NTK).