

LECTURE 19

THE KERNEL METHOD

- The proof of Grothendieck's Inequality in lecture 18 is based on the existence of maps $\phi, \psi: \mathbb{R}^d \rightarrow \mathbb{R}^N$ s.t.

$$\langle \phi(u), \psi(v) \rangle = \sin\left(\frac{\beta_n}{2} \langle u, v \rangle\right) \quad \forall u, v \in \mathbb{R}^d \text{ unit}$$

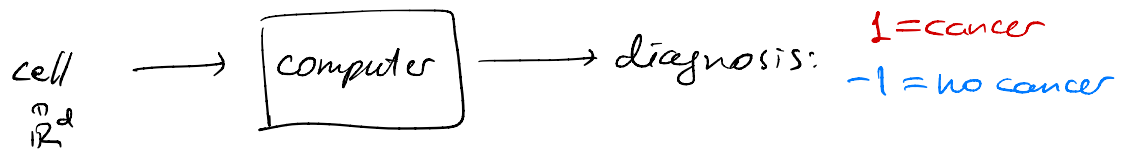
realizes nonlinearity.

- This idea leads to the kernel trick in Machine Learning (ML)

Our first example of a ML task:

BINARY CLASSIFICATION

- Want: automatic cancer diagnostics



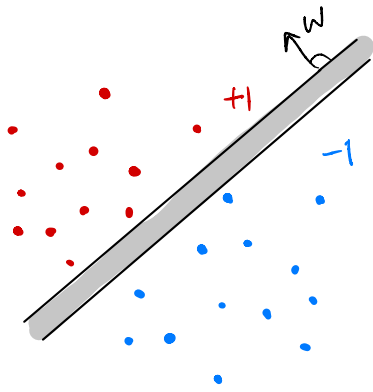
Each cell = d parameters, e.g. features of a cell
(radius, perimeter, texture, symmetry, ...) $\in \mathbb{R}^d$

$\Rightarrow$ Want $f: \mathbb{R}^d \rightarrow \{\pm 1\}$.

$$f(x) = y$$

↑ test ← prediction

- Training data: $(x_i, y_i), i=1, \dots, n$
= n cells that are already classified by doctors



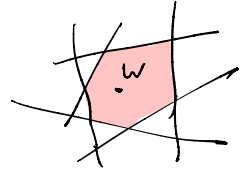
- First, assume linear separation :

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

$$\boxed{\langle w, x_i \rangle y_i > 1 \quad \forall i} \quad (*)$$

- Learning : find $w \in \mathbb{R}^d$ that satisfies $(*)$

$(*)$ = Intersection of half-spaces
 $\Rightarrow$ Linear (feasibility) program.
 Tractable.



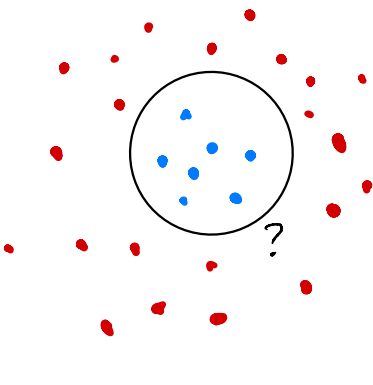
REMARK: generally, we include offset: $\langle w, x \rangle + b = 0$
 but can include it into w :
 $\langle w \oplus b, x \oplus 1 \rangle = 0$

- Prediction : $\forall$ new cell $x \in \mathbb{R}^d$, diagnose it :

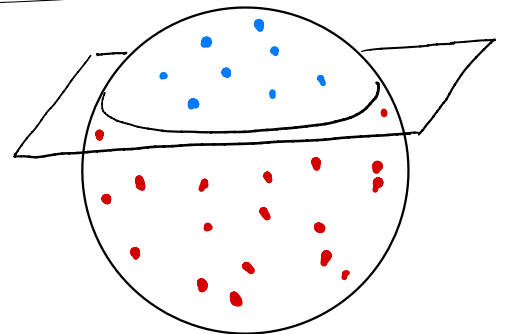
$$f(x) = \begin{cases} 1 \text{ (cancer)} & \text{if } \langle w, x \rangle > 1 \\ -1 \text{ (no cancer)} & \text{if } \langle w, x \rangle < -1 \end{cases} = \boxed{\text{sign } \langle w, x \rangle}$$

- "Support Vector Machine" (SVM), "Perceptron"

- What if there is NO linear separation? Transform :

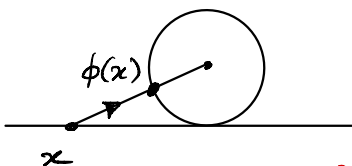


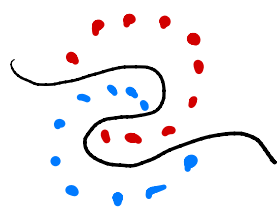
$\xrightarrow{\phi}$
 "Feature map"



e.g. stereographic projection

Transformed data is linearly separable
 $\Rightarrow$ Proceed as before.



- What if  ? How do we find ϕ ?

NO NEED :

KERNEL TRICK

- Rewrite the learning algorithm in terms of inner products $\langle x_i, x_j \rangle$:

WLOG $w = \sum_{j=1}^n \alpha_j x_j$. $\langle w, x_i \rangle = \sum_{j=1}^n \alpha_j \langle x_i, x_j \rangle$

$\left \begin{array}{l} \text{Find } (\alpha_j)_{j=1}^n : \sum_{j=1}^n y_i \alpha_j \langle x_i, x_j \rangle > 0 \quad \forall i \\ \text{Prediction: } f(x) = \text{sign} \left(\sum_{j=1}^n \alpha_j \langle x, x_j \rangle \right) \end{array} \right $ ↖ Linear program	$\xrightarrow[\text{transform}]{\phi}$	$\left \begin{array}{l} \text{Find } (\alpha_j)_{j=1}^n : \sum_{j=1}^n y_i \alpha_j \langle \phi(x_i), \phi(x_j) \rangle > 0 \quad \forall i \\ \text{Prediction: } f(x) = \text{sign} \left(\sum_{j=1}^n \alpha_j \langle \phi(x), \phi(x_j) \rangle \right) \end{array} \right $ still a linear program ↗
--	--	--

- Denote $K(x, y) := \langle \phi(x), \phi(y) \rangle$ "Kernel" $\Rightarrow$

$$\left| \begin{array}{l} \text{Find } (\alpha_j) : \sum_j y_j \alpha_j K(x_i, x_j) > 0 \quad \forall i \\ \text{Prediction: } \text{sign} \left(\sum_j \alpha_j K(x, x_j) \right) \end{array} \right| \quad (*)$$

(still a linear program).

- Forget ϕ . Just choose a suitable kernel $K(\cdot, \cdot)$ and solve (*). "Kernel SVM"

- Ex: $K(x, y) = \exp \left(- \frac{\|x - y\|^2}{2\sigma^2} \right)$ "Radial Basis Function" Kernel (RBF)
|| nontrivial

$\langle \phi(x), \phi(y) \rangle$ but we don't need to know ϕ ! Just use in (*)

$\|\phi(x)\|_2^2 = K(x, x) = \exp(0) = 1 \Rightarrow \text{maps } \mathbb{R}^d \rightarrow \underline{\text{sphere}}.$



- "Kernelizing" a ML algorithm:

- Express it in terms of inner products of data $\langle x_i, x_j \rangle$
- Replace all instances $\langle x, y \rangle \mapsto K(x, y)$ for a suitable choice of the kernel K .