

ML & Optimization

Lecture 1 8/30/2021

✓ - Logistics

- What is optimization?

- What is ML?

- Example ML problems & issue of scalability

What is Optimization

$$J: X \rightarrow \mathbb{R}$$

$$x^* = \operatorname{argmin}_{x \in X} J(x)$$

X — some set
(e.g. hypothesis space)

J — objective function

x^* — "best" element in X

n.b. we can also equivalently find x^* as the solution to a maximization problem

$$x^* = \operatorname{argmax}_{x \in X} -J(x)$$

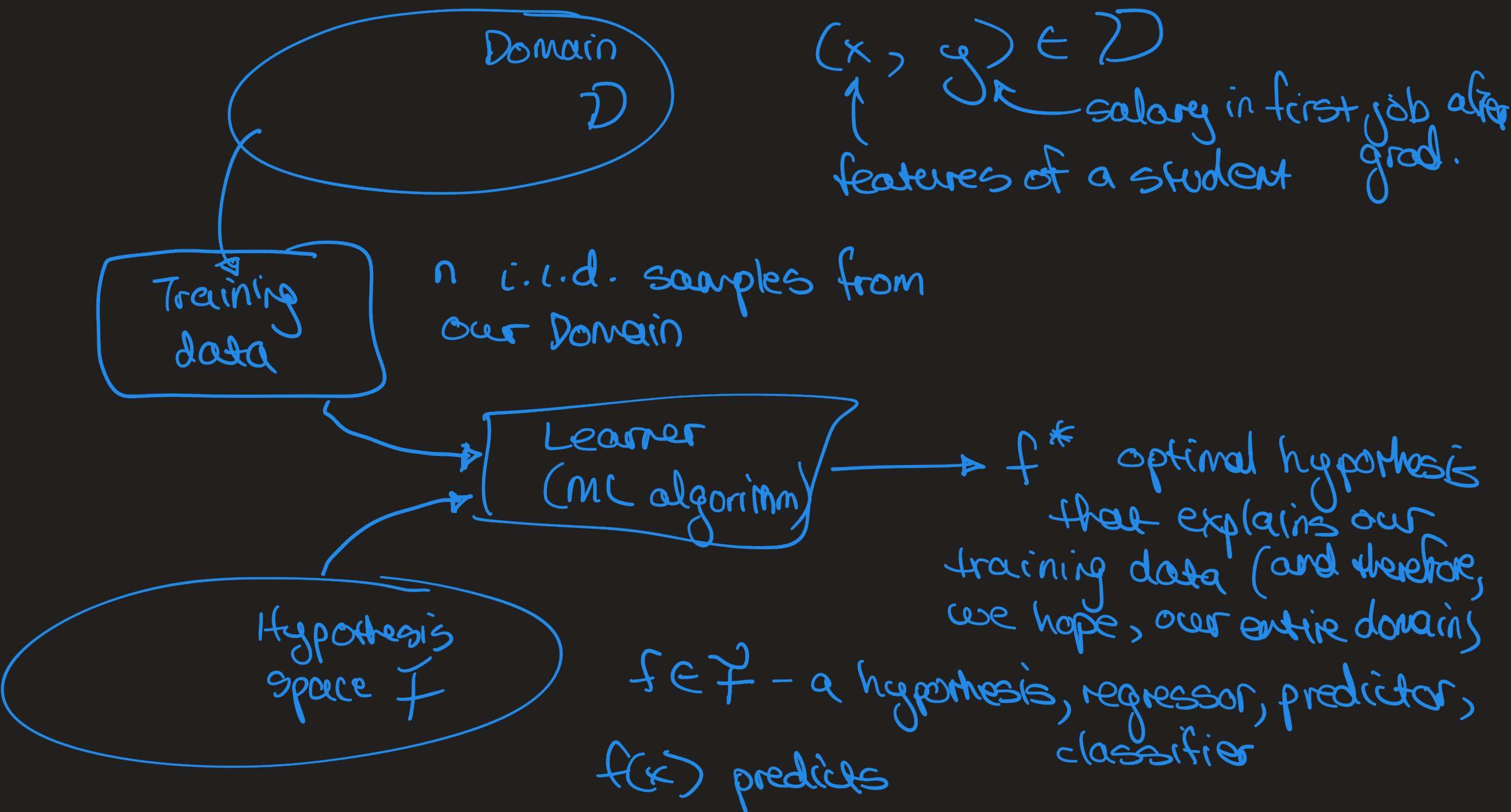
in terms of difficulty, maximization & minimization are the same

Challenges of optimization?

- constraints
- ambiguity in matching your task w/ a specific objective function
- these problems are extremely difficult when there isn't a specific structure we can exploit
- scale of the problem

What is Machine Learning?

"finding a pattern from a set of data"



What do I mean by an optimal hypothesis?

l - loss function $l(\cdot, \cdot)$

$f(x)$ - prediction of a model on an input

y - true label for that input

$$f^* = \underset{f \in \mathcal{F}}{\operatorname{argmin}} \mathbb{E}_{\mathbb{P}} l(f(x), y)$$

$$\hat{f} = \underset{f \in \mathcal{F}}{\operatorname{argmin}} \frac{1}{n} \sum_{i=1}^n l(f(x_i), y_i)$$

Examples

unsupervised: k-nearest neighbors (knn)

supervised: support vector machines (svms)

A high-level overview of the ML process:

- formulate a task
- collect appropriate data
- choose an appropriate model class for my problem
- formulate an optimization problem to fit the model on my data
- solve the optimization prob (potentially after modifying it to be easier to solve)

Ex: k-nn

Problem: given a data $\hat{p}_x$ from $\mathcal{D}$ predict the label
 $y \in \{-1, 1\}$

Data: $X = \{(x_i, y_i)\}_{i=1}^n$ where $x_i \in \mathbb{R}^d$ and $y_i \in \pm 1$

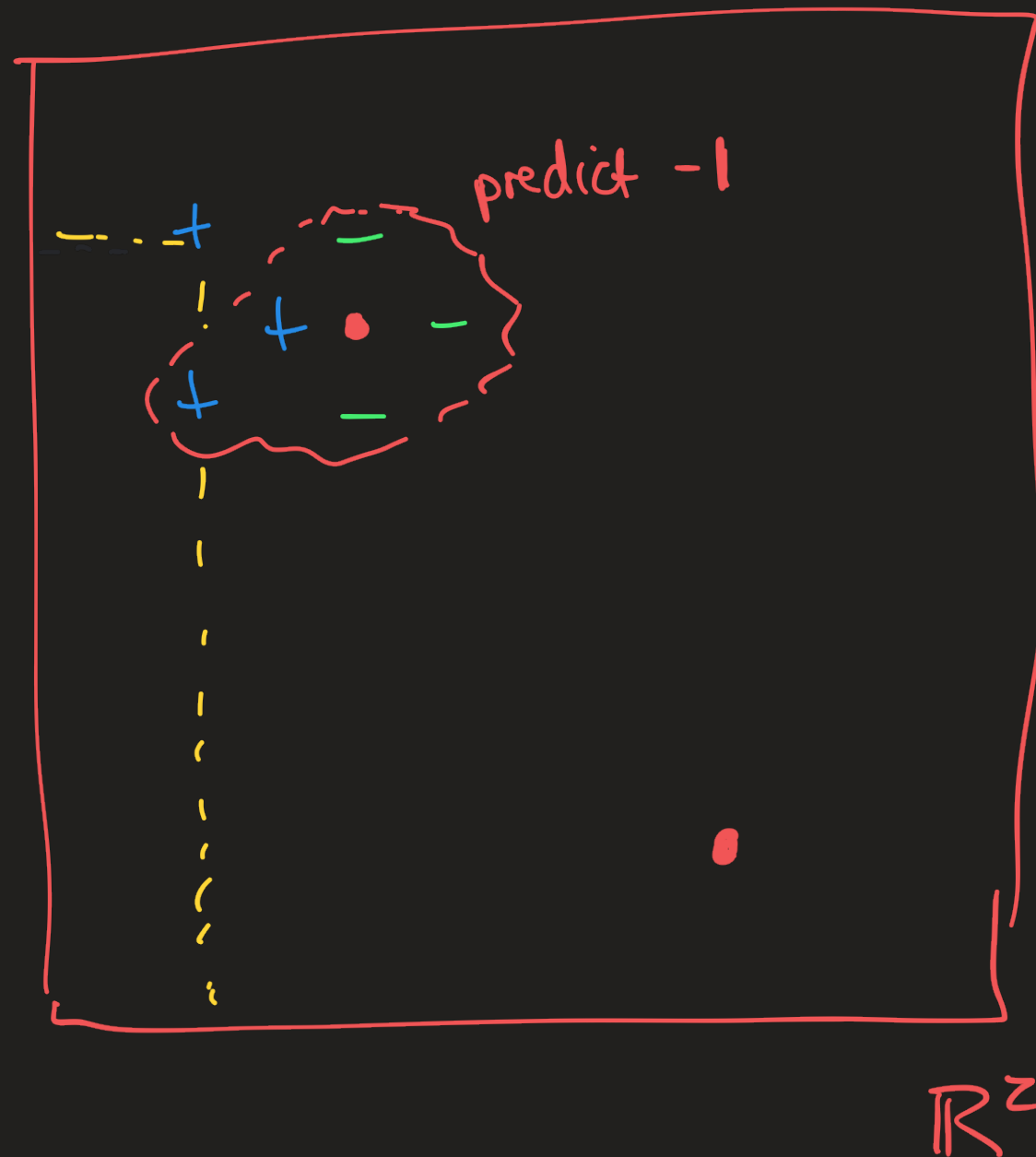
Model: k-nn

$$f(x) = \text{sign} \left(\sum_{x_j \in \mathcal{N}_k(x)} y_j \right) \quad \text{sign}(x) = \begin{cases} 1 & x > 0 \\ -1 & x < 0 \\ 0 & x = 0 \end{cases}$$

$\mathcal{N}_k(x) = \{x_j : x_j \text{ is one of the } k \text{ nearest neighbors of } x \text{ in the training data set}\}$

k-nn

$k=5$



$$O(h^d)$$

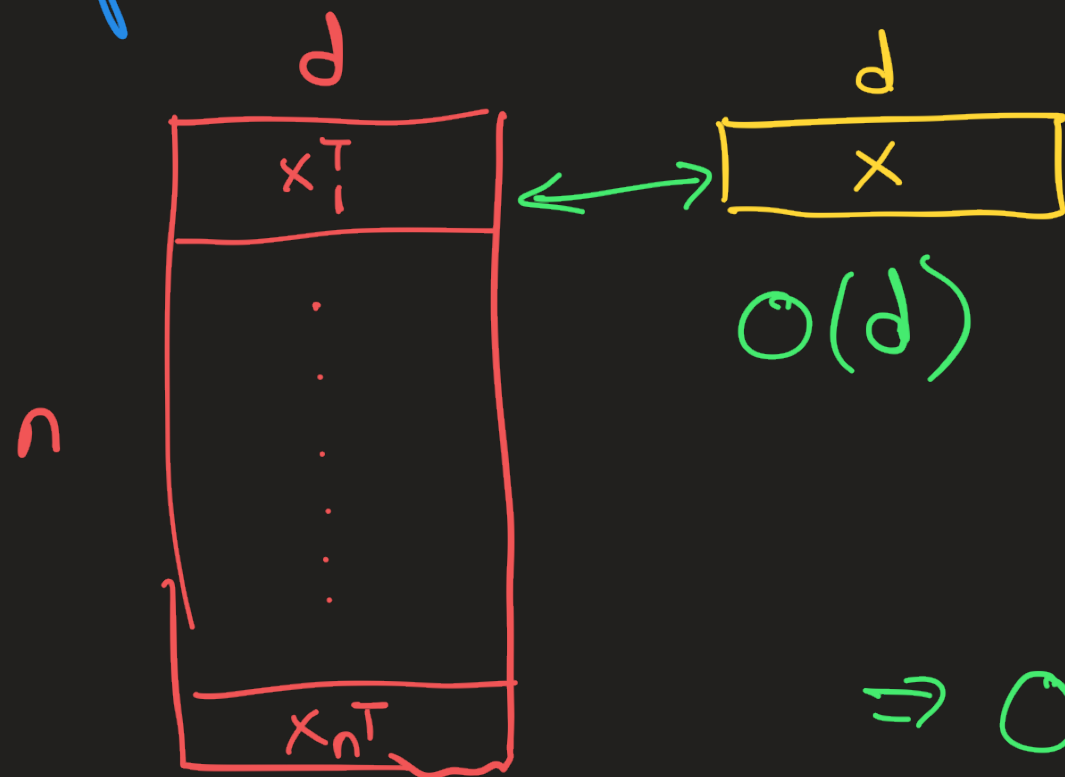
Optimization given a new point x , find $\mathcal{N}_k(x)$
and compute $f(x)$

Cost of naive k-NN

$$\underline{\underline{O(nd + n \log n + k)}}$$

very expensive
if $n \gg 1$
if $d \gg 1$

- compute distance of x from all n training data pts



$$\|x - x_i\|_2$$

$$O(d)$$

$\Rightarrow O(nd)$ operations to compute
distance of x to all training
data points

$$X \in \mathbb{R}^{n \times d}$$

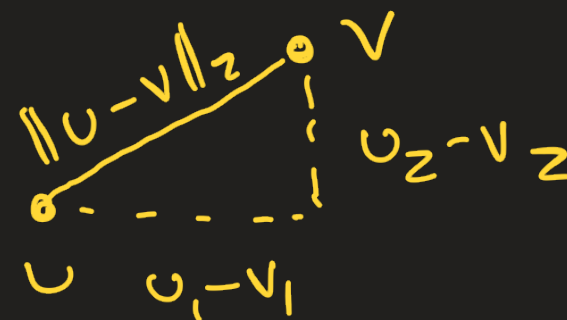
- Find the k smallest distances
to determine $\mathcal{N}_k(x)$

$$O(n \log n)$$

- Sum over $\mathcal{N}_k(x)$ to compute $f(x)$

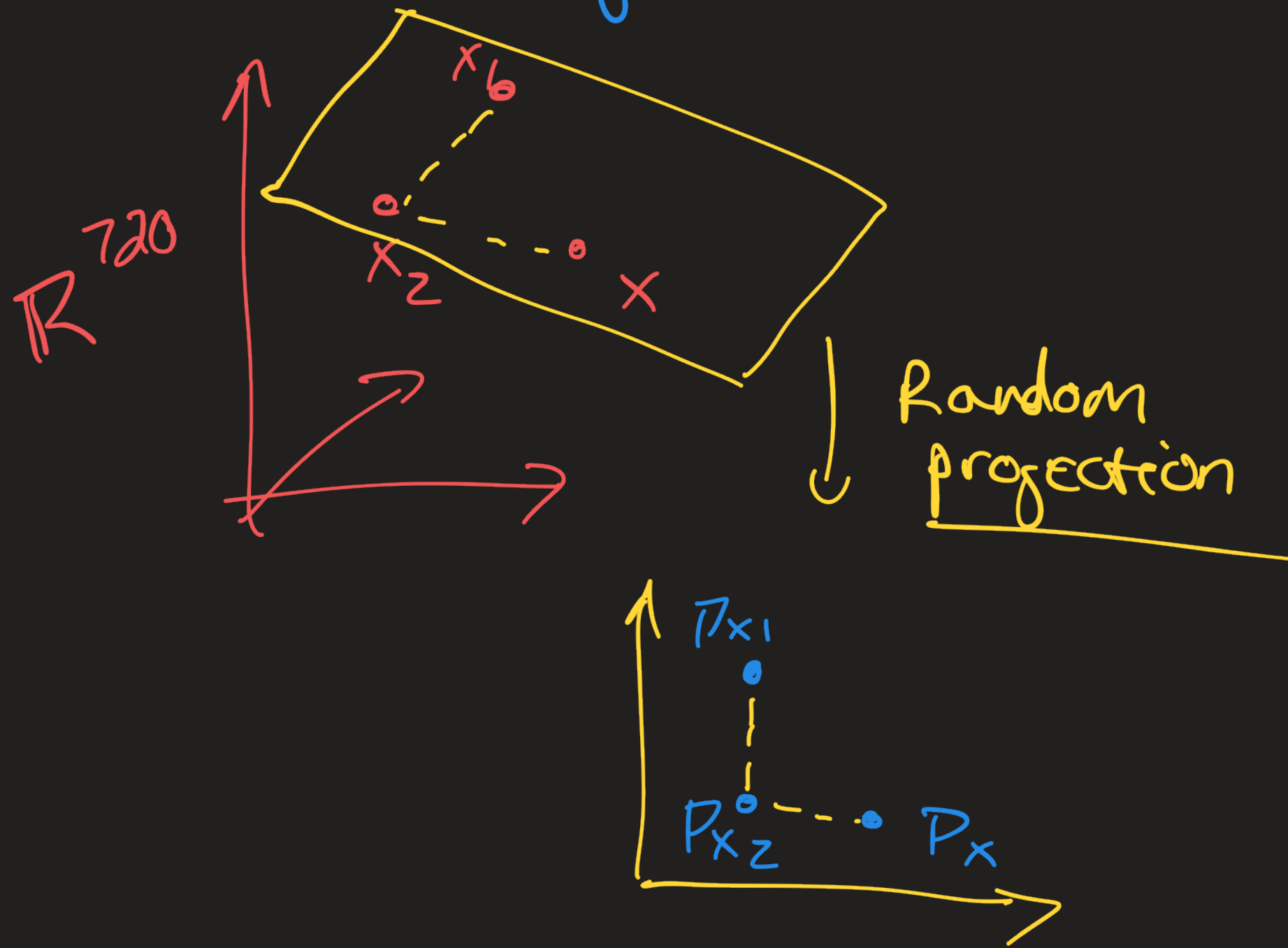
$$O(k)$$

$$\|u - v\|_2^2 = \sum_{i=1}^d (u_i - v_i)^2$$



Idea for reducing cost

We use what is called "approximate nearest neighbors"
Project our datapoints from $\mathbb{R}^d$ down to $\mathbb{R}^c$ where $c \ll d$, and then find the nearest neighbors there.



Approximate Nearest Neighbors for k-nn

- Sample a random matrix $P \in \mathbb{R}^{c \times d}$ where $c = O(\log n)$
- Compute the distances $\|Px_i - Px\|_2^2 \approx \|x_i - x\|_2^2$
- Select $N_k(x)$ using the distances $\|Px_i - Px\|_2^2$
- Compute $f(x)$

Offline component: compute P, Px_i for all ^{training} data points

$$O(cd) + O(ncd)$$

↑
computing P

↑
computing Px_i for all training pts

Online : compute P_x , compute $\|P_x - P_{x_i}\|_2$ for all training data, compute $\mathcal{N}_K(x)$, compute $f(x)$

$$\begin{array}{ccccccc} \mathcal{O}(cd) & + & \mathcal{O}(nc) & + & \mathcal{O}(n \log n) & + & \mathcal{O}(k) \\ \uparrow & & \uparrow & & \uparrow & & \uparrow \\ P_x & & \|P_x - P_{x_i}\|_2 & & \mathcal{N}_K(x) & & f(x) \end{array}$$

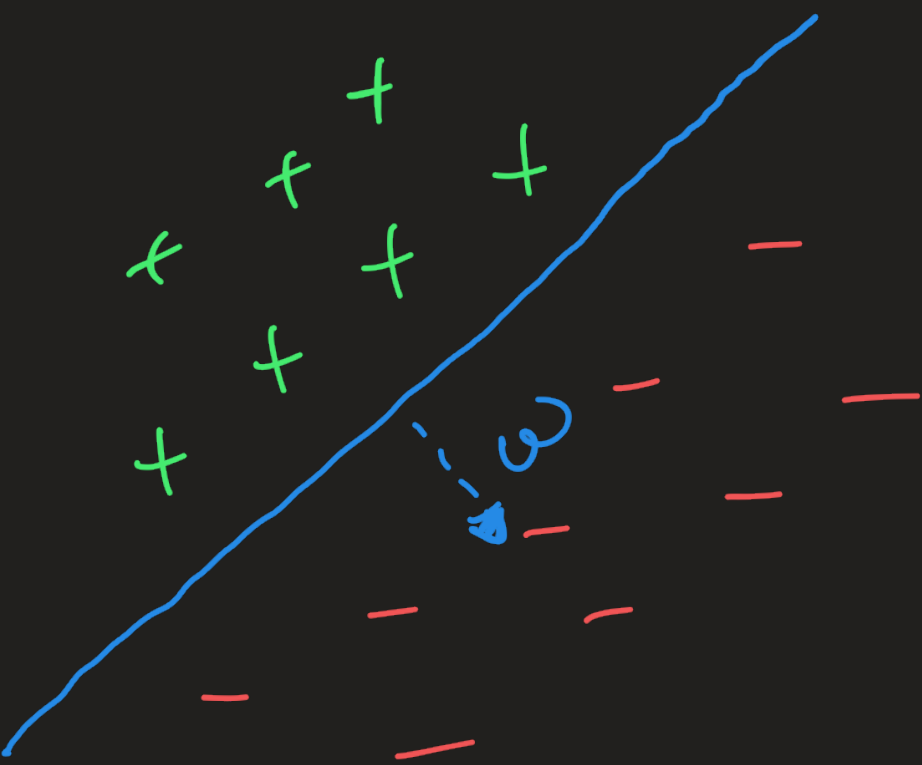
Comparison: (online portion)

Naïve knn : $\mathcal{O}(nd)$

ANN knn : $\mathcal{O}(nc)$

where $c \ll d$

Ex: Binary classification using SVMs

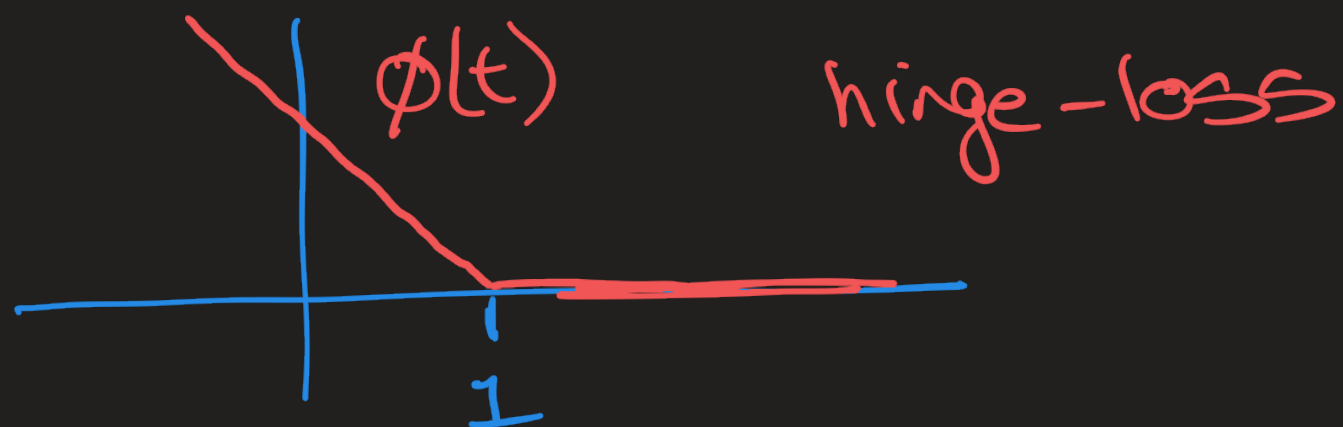


$$\mathcal{F} = \left\{ f : f(x) = \text{sign}(\omega^T x) \right. \\ \left. \text{for a } \omega \in \mathbb{R}^d \right\}$$

To choose our optimization problem for fitting an SVM

- penalize the model if y_i and $\langle \omega, x_i \rangle$ have different signs
- penalize the model if $y_i \langle \omega, x_i \rangle < 1$
- avoid trivially confident models (do not let $\|\omega\|_2 \rightarrow \infty$)

Let $\phi(t) = \max(1-t, 0) = (1-t)_+$



Let our loss function $l(f(x_i), y_i) = \phi(y_i \langle \omega, x_i \rangle)$

Two setups for the optimization prob of learning an SVM using the hinge-loss:

1) $\omega^* = \underset{\|\omega\|_2 \leq C}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n \phi(y_i \langle \omega, x_i \rangle) \quad (\text{Constrained Optim})$

2) $\omega^* = \underset{}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n \phi(y_i \langle \omega, x_i \rangle) + \lambda \|\omega\|_2^2 \quad (\text{Regularized Optim})$