

## ML and Optimization Lecture 2

- Challenges of optimization
- What is (supervised) Machine Learning
- Ex: Support Vector Machines for binary classification

# Challenges of Optimization

$$\Theta^* \in \operatorname{argmin}_{\Theta \in \mathcal{I}} J(\Theta)$$

constrained optimization

$$\Theta^* \in \operatorname{argmin} J(\Theta)$$

unconstrained optimization

- dealing w/ constraints (ex: OLS vs NNLS: Lec 1)
- ambiguity in matching the task w/ a specific objective function (ex:  $\|\cdot\|_2^2$  or  $\|\cdot\|_1$  for linear regression)
- these problems are extremely difficult to solve or even check a solution for optimality when there is not a specific structure we can exploit (ex: sparse linear regression)
- scale of the problem (ex: large-scale linear regression)

# What is Machine Learning?

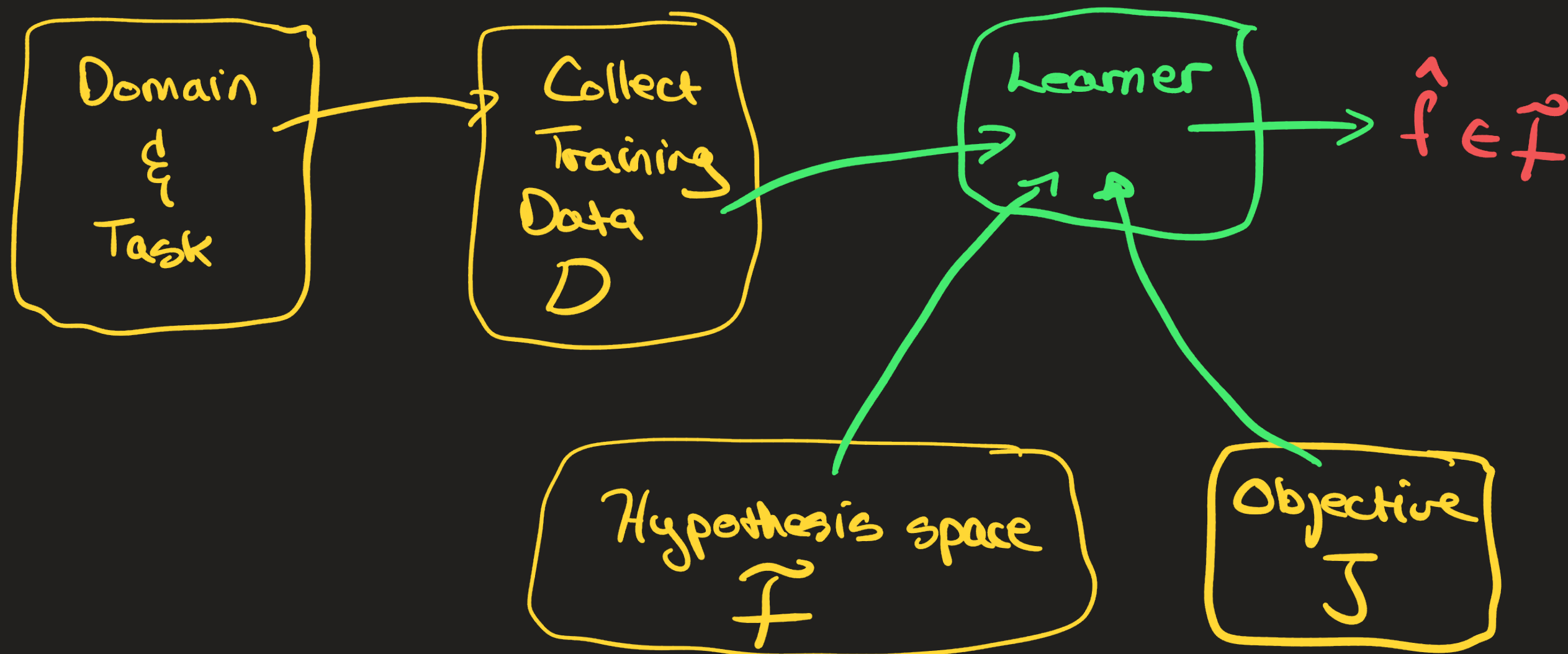
"Finding a predictable pattern from a set of data"

accurate in the presence of randomness       $f(x) = y$       randomly selected observations of the form  $(x, y)$

Assumptions that make ML possible:

- 1) there is an accurate function
- 2) that function can be learned from data

## ML Diagram



Ex: Digit recognition (MNIST)

Given  $x \in \mathbb{R}^{28 \times 28} \sim \mathbb{R}^d$  where  $d = 28 \times 28$

grayscale image of digits 0-9

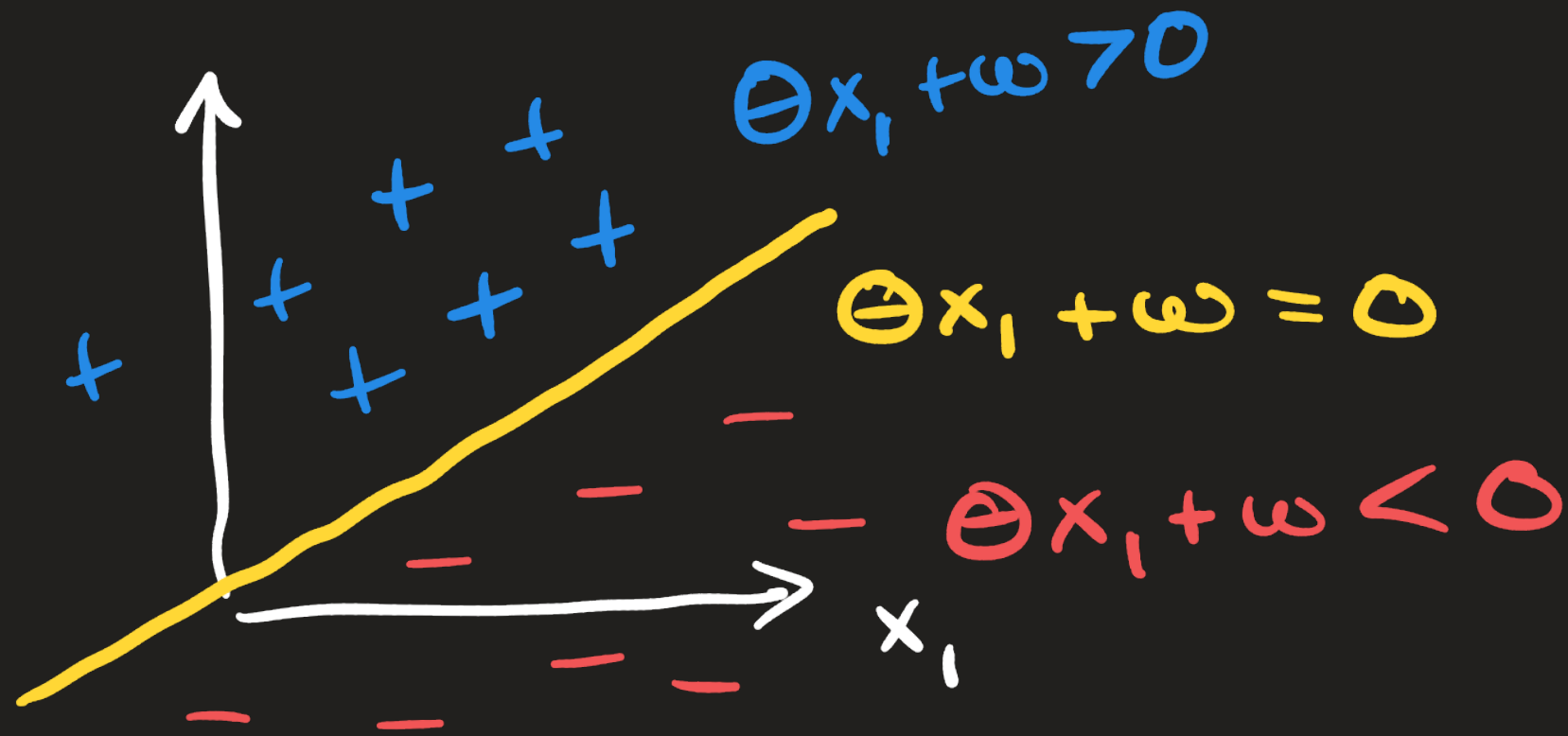
Task: predict which digit is in the image

Simplify for today to predicting if is 0 (class 1)  
or not (class -1)

Hypothesize that it suffices to consider thresholded  
affine combination so the pixels in  $x$

$$\mathcal{F} = \{ f \mid f(x) = \text{sign}(\theta^T x + \omega) \text{ for } \theta \in \mathbb{R}^d \\ \text{and } \omega \in \mathbb{R} \}$$

Geometrically, this is equivalent to assuming linear separability of the positive and negative classes



This is an ASSUMPTION and its probably not the best one



decision surface OR separating hyperplane



How to measure how good a model is on the training data?

Maybe we can choose  $\Theta, \omega$  to maximize accuracy

$$\hat{\Theta}, \hat{\omega} = \operatorname{argmax} J(\Theta, \omega) = \frac{1}{n} \sum_{i=1}^n \operatorname{sign}(\Theta^T x_i + \omega) y_i$$

But optimizing 0-1 accuracy is difficult:

$J$  is discontinuous w.r.t.  $\Theta$  and  $\omega$

We want a nice smooth objective. Idea: penalize models that make predictions close to the decision surface. Notice that if  $\operatorname{sign}(\Theta^T x_i + \omega) = y_i$  and  $|\Theta^T x_i + \omega|$  is large then the model is confident

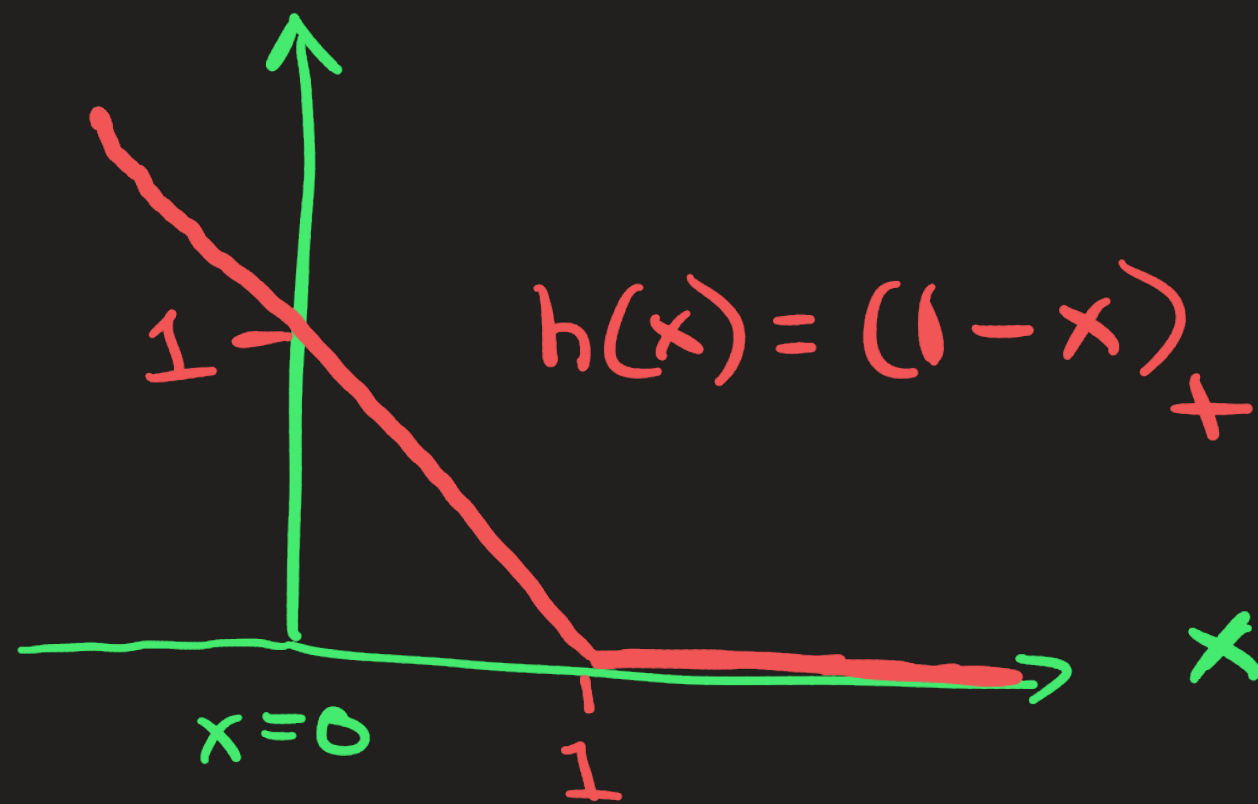
Notice that  $|\Theta^T x_i + \omega| = (\Theta^T x_i + \omega) y_i$   
when the model is correct

Take

$$J(\Theta, \omega) = \frac{1}{n} \sum_{i=1}^n h((\Theta^T x_i + \omega) y_i)$$

where  $h$  is the "hinge loss"

$$h(x) = (1 - x)_+ = \begin{cases} 0 & \text{if } x \geq 1 \\ 1 - x & \text{otherwise} \end{cases}$$





With this choice we choose our best model for binary classification by solving

$$(*) \quad \hat{\Theta}, \hat{\omega} = \underset{\Theta, \omega}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n (1 - (\Theta^T x_i + \omega) y_i)_+$$

This objective is continuous and a.e. differentiable. One last issue to address in our choice of objective: if  $\alpha > 0$ , then the two models  $(\Theta, \omega)$  and  $(\alpha\Theta, \alpha\omega)$  define the same decision surface:

$$\Theta^T x + \omega = 0 \iff (\alpha\Theta)^T x + \alpha\omega = 0$$

so  $(*)$  is ill-posed: there are infinitely many equally well-performing models

We break this indeterminacy by using regularization so our final objective is

$$\hat{\Theta}, \hat{\omega} = \underset{\Theta, \omega}{\operatorname{argmin}} \underbrace{\frac{1}{n} \sum_{i=1}^n (1 - (\Theta^T x_i + \omega) y_i)_+}_{\text{data-fitting term}} + \underbrace{\lambda \|\Theta\|_2^2}_{\text{regularization term}}$$

Here  $\lambda > 0$  is called the regularization parameter and  $\|\Theta\|_2^2$  is called  $\ell_2$ -regularizer or weight decay

By varying  $\lambda$  (a hyperparameter) you can prioritize fitting the data more or having a less complex models

This general class of models for binary classification w/ targets in  $\pm 1$  is called  
Support Vector Machines (SVM)

$$\mathcal{F}_{\text{svm}} = \{f \mid f(x) = \text{sign}(\Theta^T x + \omega)\}$$

This particular objective is called soft-margin SVM.

$$J(\Theta, \omega) = \frac{1}{n} \sum_{i=1}^n (1 - (\Theta^T x_i + \omega) y_i)_+$$

# Generalize to supervised ML

$\ell$  — loss function  $\ell(\hat{y}, y) \in \mathbb{R}$   
measures the loss in predicting  $\hat{y}$  when the true value is  $y$

$f(x)$  — prediction of a candidate model  $f$  for input  $x$

$\mathcal{D} = \{(x_i, y_i)\}_{i=1}^n$  — training data set

$f^* = \operatorname{argmin}_{f \in \mathcal{F}} \underbrace{\mathbb{E}_{(x,y) \sim \mathcal{D}} \ell(f(x), y)}_{\text{population risk}}$

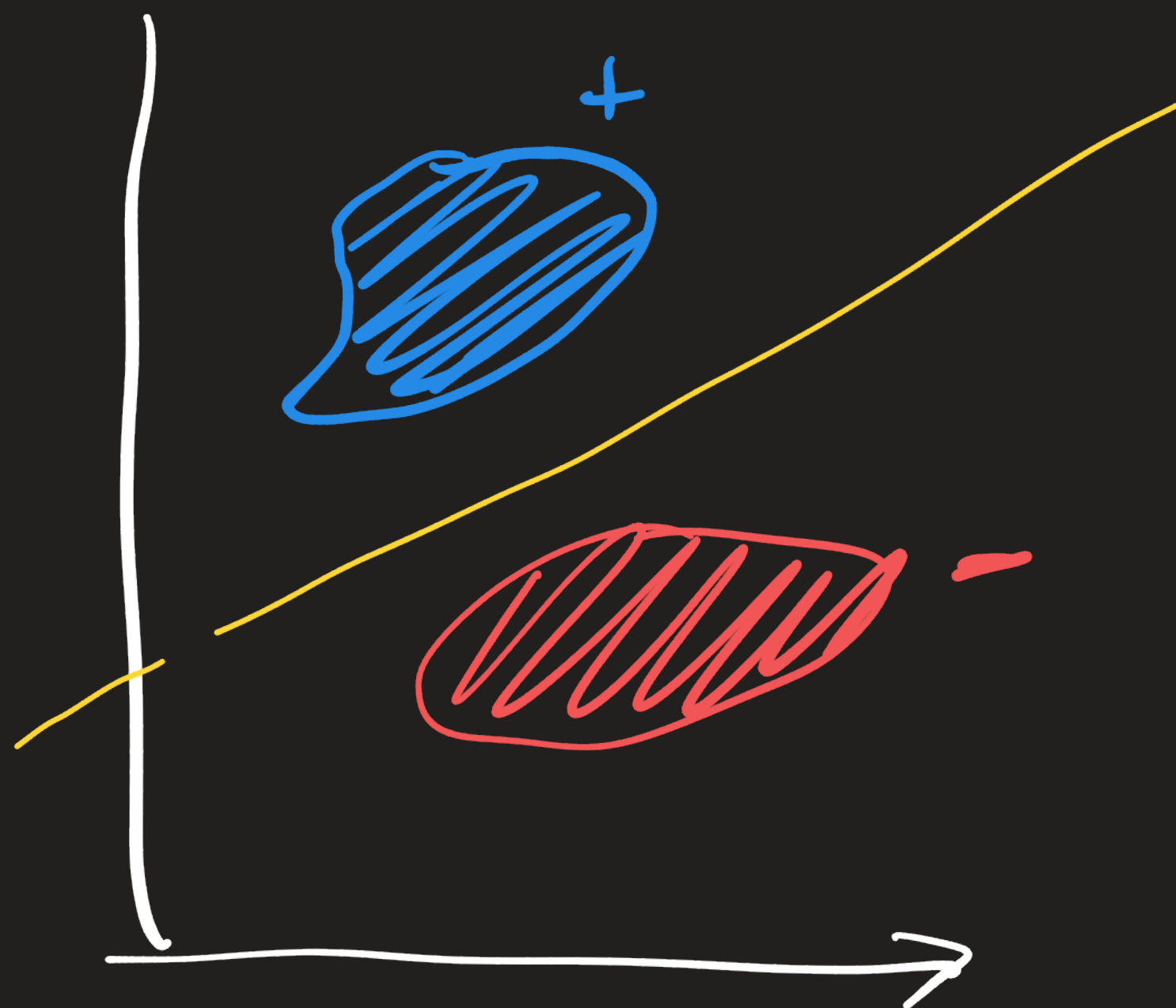
In this class we use Empirical Risk minimization.

- Assume  $\mathbb{E}_{(x,y) \sim \mathcal{P}} \ell(f(x), y) \approx \frac{1}{n} \sum_{(x_i, y_i) \in \mathcal{D}} \ell(f(x_i), y_i)$

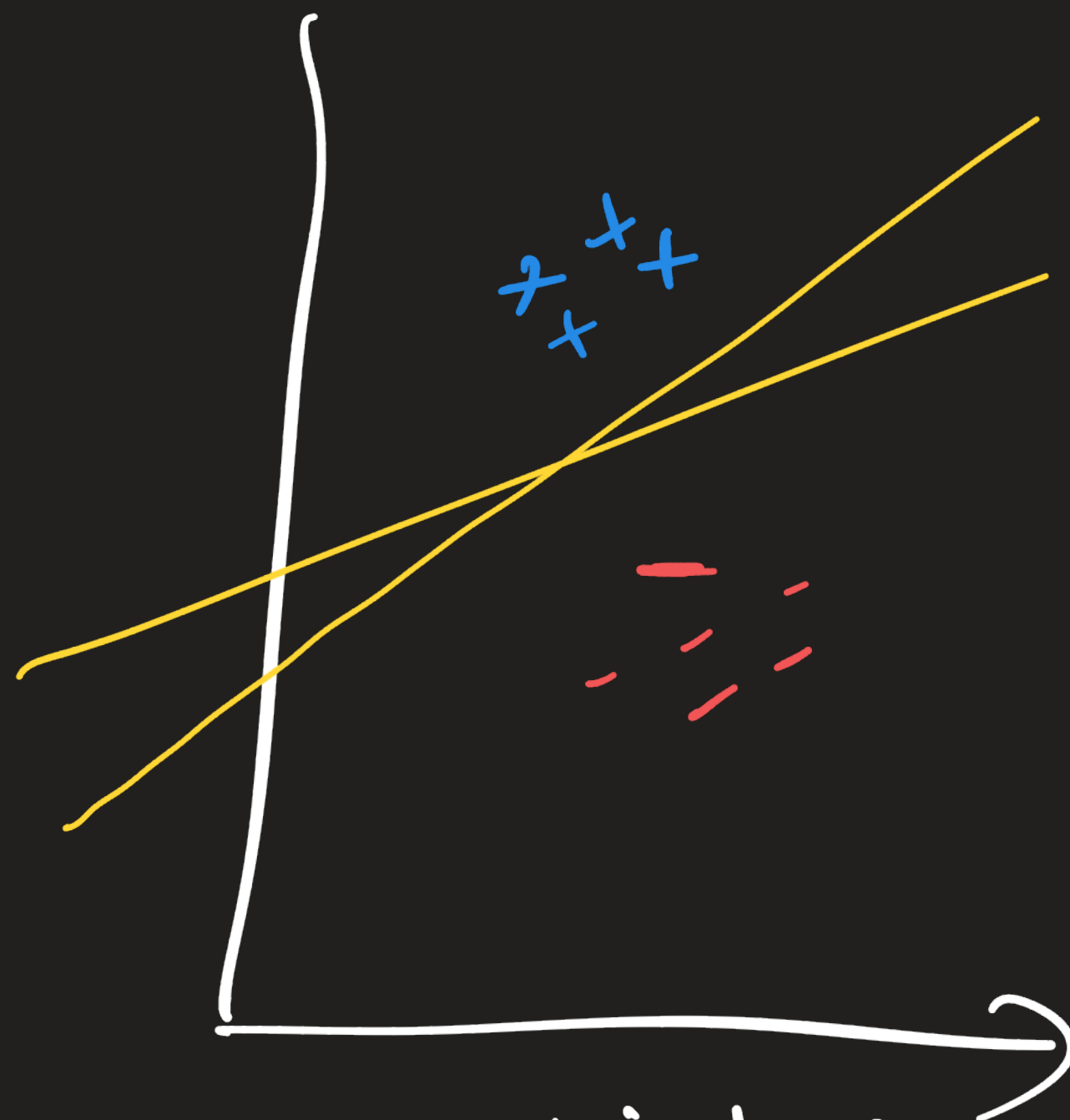
and solve

$$\hat{f} = \underset{f \in \mathcal{F}}{\operatorname{argmin}} \underbrace{\frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i)}_{\text{empirical risk}}$$





population risk  
minimization



empirical risk  
minimization

# High-level overview of ML

- formulate a task & pick a domain ( $\mathcal{P}$ )
- collect appropriate training data for ERM  
( $\mathcal{D}$  i.i.d. samples from  $\mathcal{P}$ )
- choose a good model class for the problem  
(SVM, logistic regression, multiclass LR,  
linear regression, kernel versions of these, DL)
- pick a loss function appropriate to the task
- solve the ERM optimization problem

