

ML and Optimization Lecture 2 9/1/2023

- What is ML?
- Example of ML: support vector machines
- Probability Theory*

* next time

what is ML?

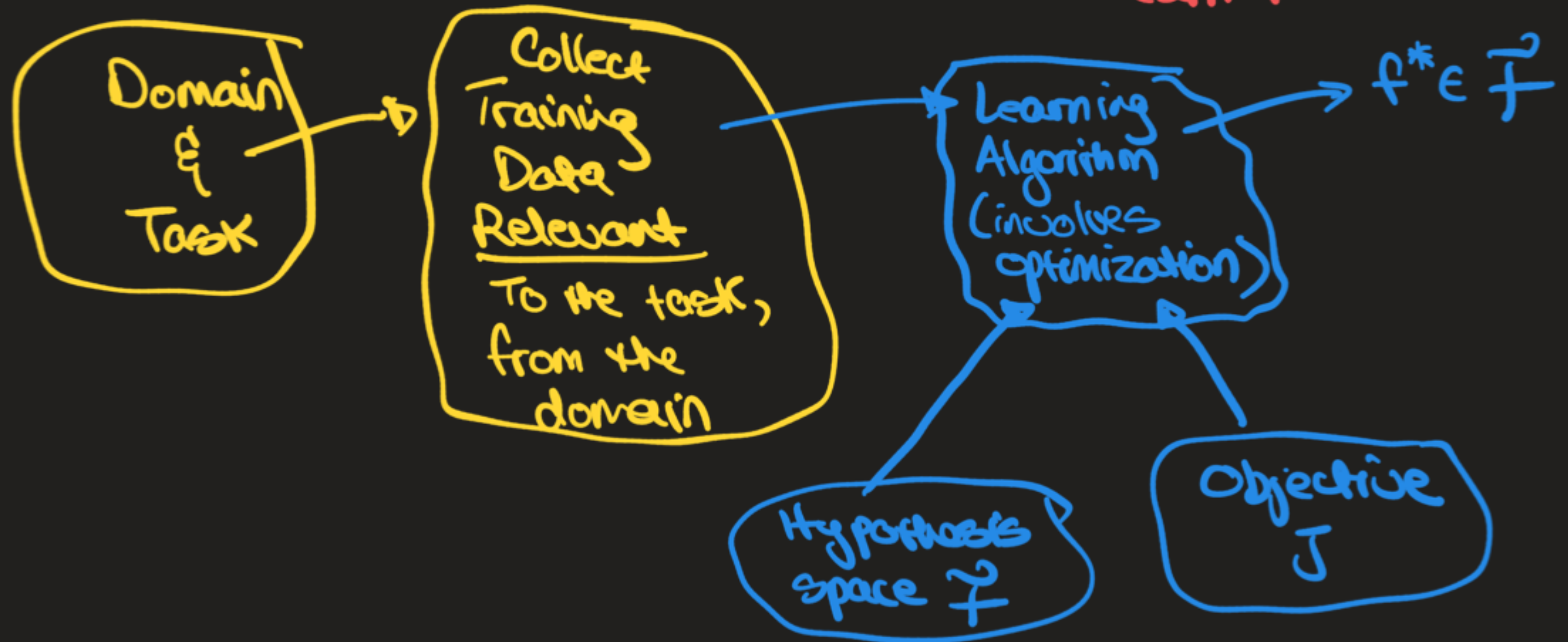
(supervised learning)

"finding a predictable pattern from a set of data"

accurate

$f: x \mapsto y$

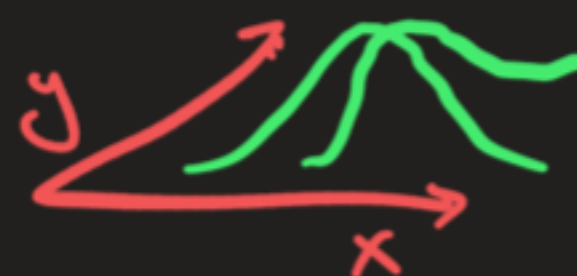
observations of
 (x, y) used to
learn f



What does it mean for f to be optimal?

First : f needs to minimize some error over the domain.

We describe domains as sets of (x, y) with a probability measure, i.e. a function that maps subsets of my domain to numbers in $[0, 1]$
write $(x, y) \sim P$



Now define/choose

l - loss function $l(\cdot, \cdot)$ measures discrepancy between prediction and true value
this depends on the task

f - function in our hypothesis space $\mathcal{F}$
(e.g. the linear functions, quadratic functions, NNs)
this depends on the task & amount of training data

Formalize ML as Empirical Risk Minimization

$$f^* = \operatorname{argmin}_{f \in \mathcal{F}} \mathbb{E}_{(x,y) \sim \mathcal{P}} \ell(f(x), y) \quad \text{Population Risk}$$

by the law
of large
numbers,
as $n \rightarrow \infty$

$$\int \ell(f(x), y) p(x, y) dx dy$$
$$\downarrow$$
$$\operatorname{argmin}_{f \in \mathcal{F}} \frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i) \quad \text{Empirical Risk}$$

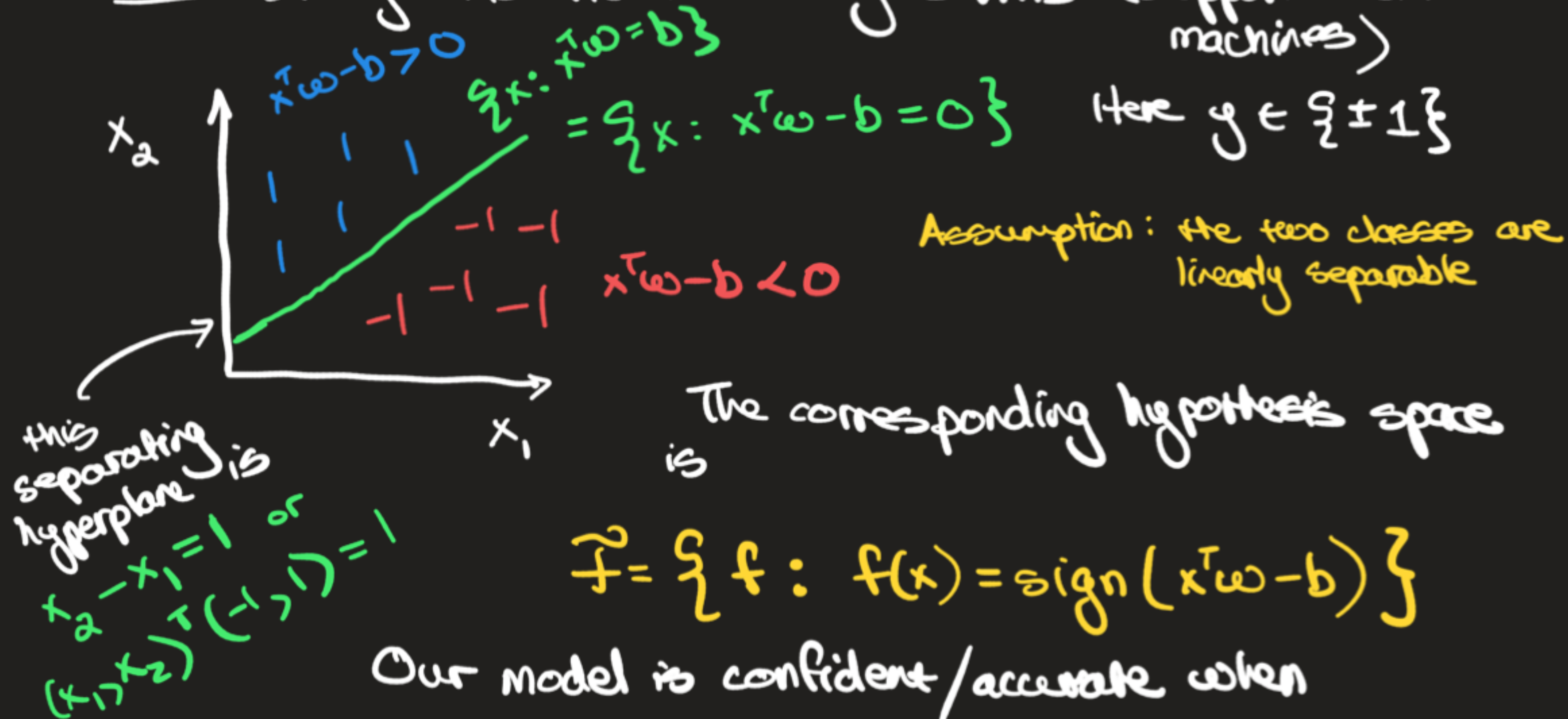
where $(x_i, y_i) \sim \mathcal{P}$ are independent,
identically distributed samples from
our domain

Empirical Risk Minimization (ERM) consists of
choosing f by minimizing the empirical risk

High-level workflow for ML (supervised ML)

- pick a domain & formulate a task (choose an appropriate loss function)
- collect i.i.d. samples from our domain
- choose an appropriate hypothesis class $\mathcal{F}$
- solve the ERM
- iterate

Ex: Binary classification using SVMs (support vector machines)



$$\mathcal{F} = \{f: f(x) = \text{sign}(x^T w - b)\}$$

Our model is confident/accurate when

$$y(x^T w - b) \gg 1$$

- the prediction $f(x) = \text{sign}(x^T w - b)$ must equal y
- the model must be confident: x is far from the separating hyperplane

Models with low confidence can be made arbitrarily highly confident: if

$$\min_i y_i (x_i^T \omega - b) = d > 0$$

then

$$\min_i y_i (x_i^T (\alpha \omega) - \alpha b) = \alpha d > 0$$

for any $\alpha > 0$

To avoid this, we restrict the norm of (ω, b)

We choose our ERM to satisfy these desiderata by:

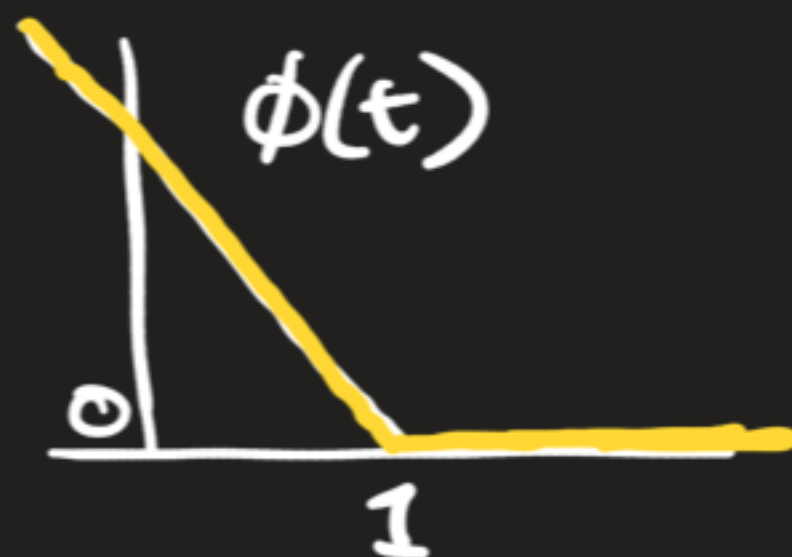
- 1) choose a loss function to ensure confidence
- 2) restrict our hypothesis space to ensure (ω, b) is bounded

Our hypothesis space becomes

$$\mathcal{F} = \{f : f(x) = \text{sign}(\omega^T x + b), \text{ where } \|\omega\|_2 \leq C\}$$

For our loss function, we use the hinge loss:

$$\phi(t) = (1 - t)_+ = \begin{cases} 1 - t & t < 1 \\ 0 & \text{otherwise} \end{cases}$$



$$\ell(y, f(x)) = \phi(y(\omega^T x - b))$$

The final ERM for the soft-margin SVM is

$$(C) \quad \min_{\substack{\omega \in \mathbb{R}^d \\ b \in \mathbb{R} \\ \|\omega\|_2 \leq C}} \quad \frac{1}{n} \sum_{i=1}^n (1 - y_i(x_i^T \omega - b))_+ \quad \text{constrained ERM}$$

or

$$(L) \quad \min_{\substack{\omega \in \mathbb{R}^d \\ b \in \mathbb{R}}} \quad \frac{1}{n} \sum_{i=1}^n (1 - y_i(x_i^T \omega - b))_+ + \lambda \|\omega\|_2^2$$

regularization coefficient
Lagrangian multiplier
or regularized version

These are equivalent: for any C there is a corresponding λ so the solutions to (C) and (L) are the same.

