

Example (regression $y \in \mathbb{R}$)

$$l(u, v) = \frac{1}{2}(u - v)^2 \quad \text{square loss / } l_2$$

$$l(u, v) = |u - v| \quad \text{absolute loss / } l_1$$

$$l(u, v) = |u - v|^p \quad \text{p-loss / } l_p$$

If we take $l = l_2$ loss, then

$$R(f) = \mathbb{E} l(f(x), y) = \mathbb{E} (f(x) - y)^2$$

mean-squared prediction error
of f (MSE)

$$= \mathbb{E}_X \left[\mathbb{E}_Y [(f(x) - y)^2 | X] \right]$$

Now that

$$\begin{aligned} & \mathbb{E}_\varphi [(f(x) - \varphi)^2 | x] \\ &= \mathbb{E}_\varphi [f(x)^2 - 2f(x)\varphi + \varphi^2 | x] \\ &= f(x)^2 - 2f(x) \underbrace{\mathbb{E}[\varphi | x]}_{\text{constant}} \\ &\quad + \underbrace{\mathbb{E}_\varphi [\varphi^2 | x]}_{\text{constant}} \end{aligned}$$

so this minimized when

$$\begin{aligned} 2f(x) - 2\mathbb{E}[\varphi | x] &= 0 \\ \Rightarrow f(x) &= \mathbb{E}[\varphi | x] \end{aligned}$$

so the predictor f^* that minimizes the MSE risk in estimating Y is the conditional expectation of Y given X , $E[Y|X]$

f^* can be arbitrarily complicated, and we restrict ourselves to a particular function class, so

$$f_{\mathcal{F}}^* = \underset{f \in \mathcal{F}}{\operatorname{argmin}} R(f)$$

can be thought of as the best estimate of $E[Y|X]$ in $\mathcal{F}$

Fact: if $d(u, v) = |u - v|$

then $f^* = \text{Med}(q(x))$

Recall we don't have access to $\mathbb{P}$,
so we do ERM

$$\hat{f}_Z = \operatorname{argmin}_{f \in \mathcal{F}} R_n(f)$$

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

and we hope

$$R(\hat{f}_Z) \approx R(f_Z^*) \approx R(f^*)$$

↑
optimization
& generalization
error

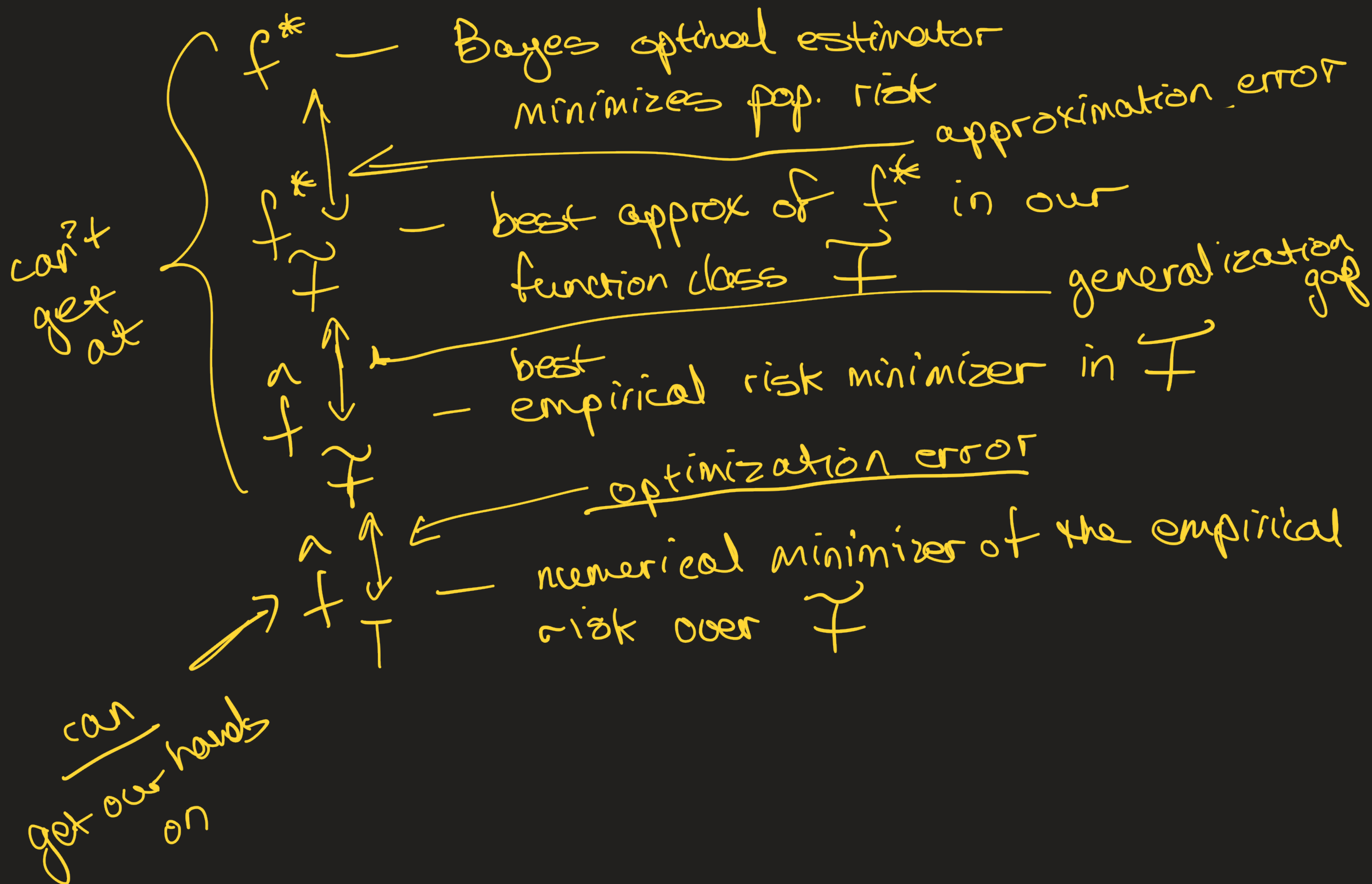
↑
approximation error:
how good is $\mathcal{F}$

An additional concern is that we use numerical optimization, so we don't actually find $\hat{f}_{\mathcal{Z}}$, instead we find $\hat{f}_T$ say after T iterations of algorithm A .

Thus the error (risk) of models learned using ERM breaks down as

$$\mathcal{E}(A, n, \mathcal{Z}, T) = \underline{R(\hat{f}_T)}$$

$$= \underbrace{R(\hat{f}_T) - R(\hat{f}_{\mathcal{Z}})}_{\text{optimiz error}} + \underbrace{R(\hat{f}_{\mathcal{Z}}) - R(f_{\mathcal{Z}}^*)}_{\text{generalization gap}} + \underbrace{R(f_{\mathcal{Z}}^*) - R(f^*)}_{\text{approximation error} + R(f^*)}$$



Ex logistic regression (Maximum Likelihood Estimation give rise to ERM)

Assume we have predictor variables $x \in \mathbb{R}^d$ and a target $y \in \{\pm 1\}$, and we can model the log odds of $y=1$ linearly in x

$$\ln\left(\frac{p}{1-p}\right) = \underline{w_0} + \sum_{i=1}^d \underline{w_i x_i}$$

interpretation: probability of success goes up if $w_i > 0$ and x_i increases
goes down if $w_i < 0$ and x_i increases

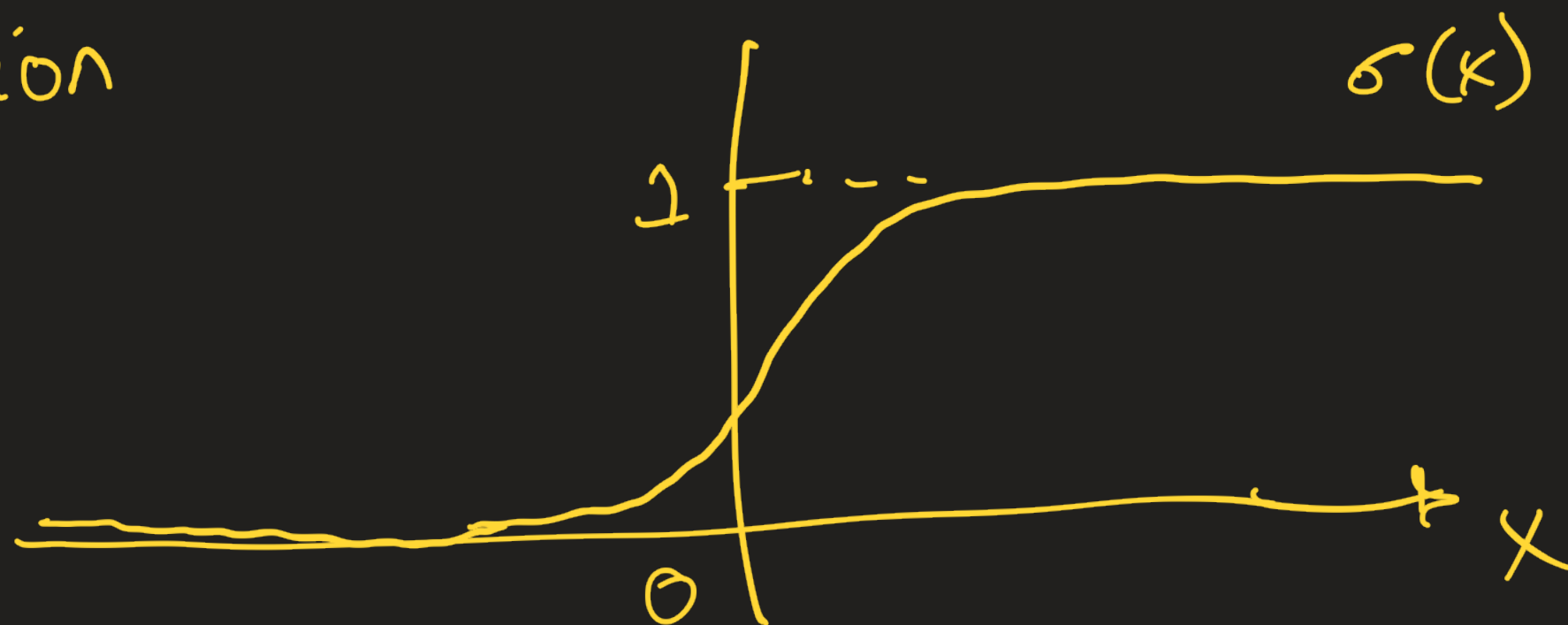
$$\Rightarrow \frac{p}{1-p} = e^{\omega_0 + \omega^T x} \Rightarrow p = \frac{e^{\omega_0 + \omega^T x}}{1 + e^{\omega_0 + \omega^T x}}$$

$$\Rightarrow p = \frac{1}{1 + e^{-\omega_0 + \omega^T x}}$$

$$p = \sigma(\omega_0 + \omega^T x)$$

where $\sigma(x) = \frac{1}{1 + e^{-x}}$ is called the sigmoid

function



Now that means

$$y \sim \text{Bern}(\sigma(\omega_0 + \omega^T x))$$

Goal: recover ω_0 & ω given training data
so in the future when we observe x
we can predict y .

Use maximum likelihood principle: choose the
 ω_0, ω that maximize the prob. of your
training data.

$$\begin{aligned}
w_0^*, w^* &= \underset{w_0, w}{\operatorname{argmax}} p(y_1, \dots, y_n) \quad \leftarrow \text{prob. of my observed data if they were drawn i.i.d. from the model} \\
&= \underset{w_0, w}{\operatorname{argmax}} p(y_1, \dots, y_n)^{1/n} \\
&= \underset{w_0, w}{\operatorname{argmax}} \left[\prod_{i=1}^n p(y_i) \right]^{1/n} \\
&= \underset{w_0, w}{\operatorname{argmax}} \log(\dots) \\
&= \underset{w_0, w}{\operatorname{argmax}} \sum_{i=1}^n \frac{1}{n} \log(p(y_i)) \\
&= \underset{w_0, w}{\operatorname{argmin}} \frac{1}{n} \sum_{i=1}^n -\log(p(y_i))
\end{aligned}$$

$y \sim \text{Bern}(\sigma(w_0 + w^T x))$

The quantity
$$-\log p(y_1, \dots, y_n)$$
$$= \sum_{i=1}^n -\log(p(y_i))$$

is called the negative log-likelihood of
the model parameters w_0, w

ML principle $\Rightarrow$ choose your model params
by minimizing the negative log-likelihood

$$w_0^*, w^* = \underset{w_0, w}{\operatorname{argmin}} \quad -\frac{1}{n} \sum_{i=1}^n \log(p(y_i))$$

$$p(y_i) = \begin{cases} \sigma(w_0 + w^T x_i) & \text{if } y_i = 1 \\ 1 - \sigma(w_0 + w^T x_i) & \text{if } y_i = -1 \end{cases}$$

$$\text{fact: } \underbrace{\frac{1}{1+e^{-x}}}_{\sigma(x)} + \underbrace{\frac{e^{-x}}{1+e^{-x}}}_{\frac{1}{1+e^x}} = 1 \Rightarrow \boxed{1 - \sigma(x) = \sigma(-x)}$$

$$\Rightarrow p(y_i) = \begin{cases} \sigma(w_0 + w^T x_i) & \text{if } y_i = 1 \\ \sigma(-(w_0 + w^T x_i)) & \text{if } y_i = -1 \end{cases}$$

$$\Rightarrow p(y_i) = \sigma(y_i(\omega_0 + \omega^T x_i))$$

$$\Rightarrow \omega_0^*, \omega^* = \operatorname{argmin} -\frac{1}{n} \sum \log(\sigma(y_i(\omega_0 + \omega^T x_i)))$$

$$-\log(\sigma(x)) = \log\left(\frac{1}{\sigma(x)}\right) = \underbrace{\log(1 + e^{-x})}_{\text{log loss!}}$$

$$\Rightarrow \omega_0^*, \omega^* = \operatorname{argmin} \frac{1}{n} \sum_{i=1}^n \log(1 + e^{-y_i(\omega_0 + \omega^T x_i)})$$

logistic regression

This is ERM w/ loss

$$\ell(u, v) = \log(1 + e^{-uv})$$

Note we took

$$y \sim \text{Bern}(\sigma(\omega_0 + w^T x))$$

we could say

$$y \sim \text{Bern}(\sigma(\underline{f(x)}))$$

for $f \in \tilde{\mathcal{F}}$

then we get non-linear logistic regression

$$f^* = \arg \min_{f \in \tilde{\mathcal{F}}} \frac{1}{n} \sum_{i=1}^n \log(1 + e^{-y_i \sigma(f(x_i))})$$

Next time:

overfitting
regularization
convex optimiz

(poly regression)
(lasso, ridge regression)