

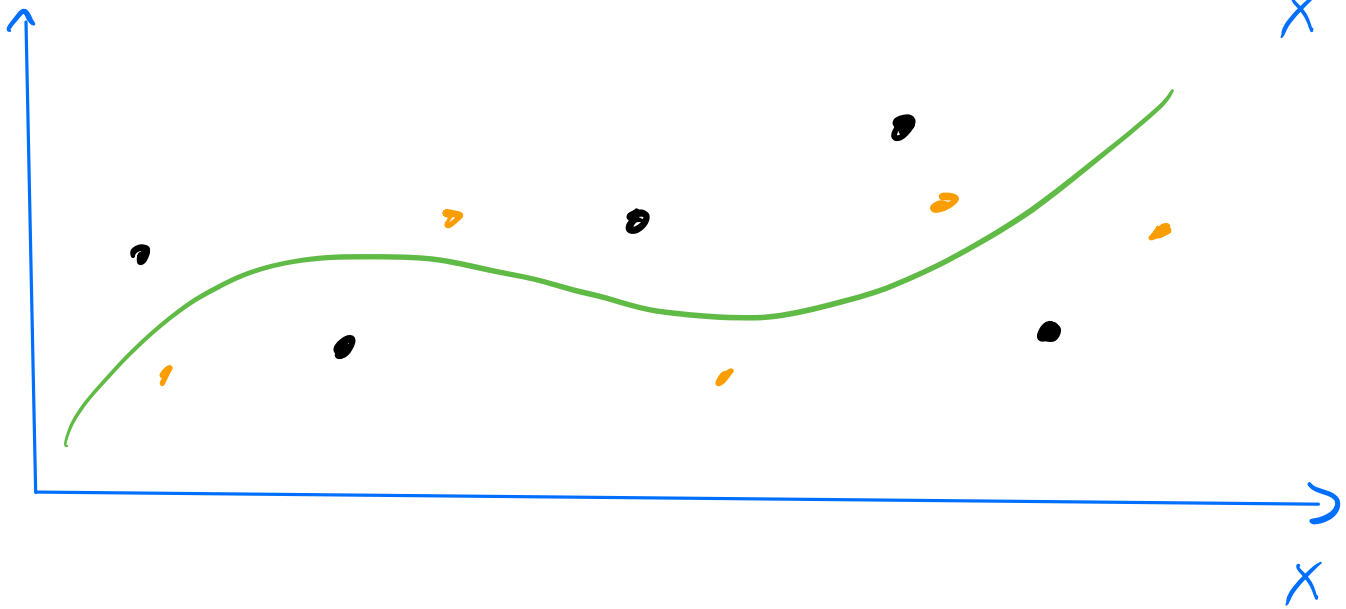
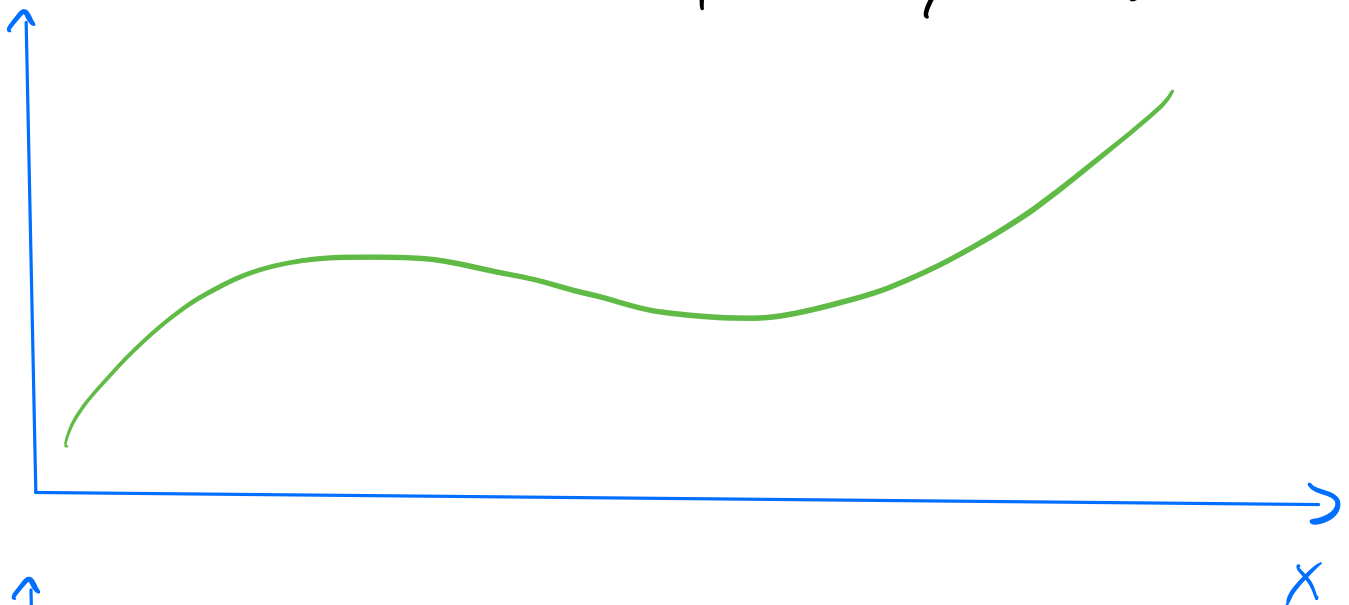
Evaluating Predictors / Models

$$\mathcal{A}(D) = \hat{f}_D = \arg \min_{f \in \mathcal{F}} \hat{L}(f) \quad D \text{ is a r.v.}$$

$$\hat{L}(f) = \frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i), \quad L(f) = \mathbb{E}[\ell(f(X), Y)]$$

$$\hat{L}(\hat{f}_D) = \frac{1}{n} \sum_{i=1}^n \ell(\hat{f}_D(x_i), y_i)$$

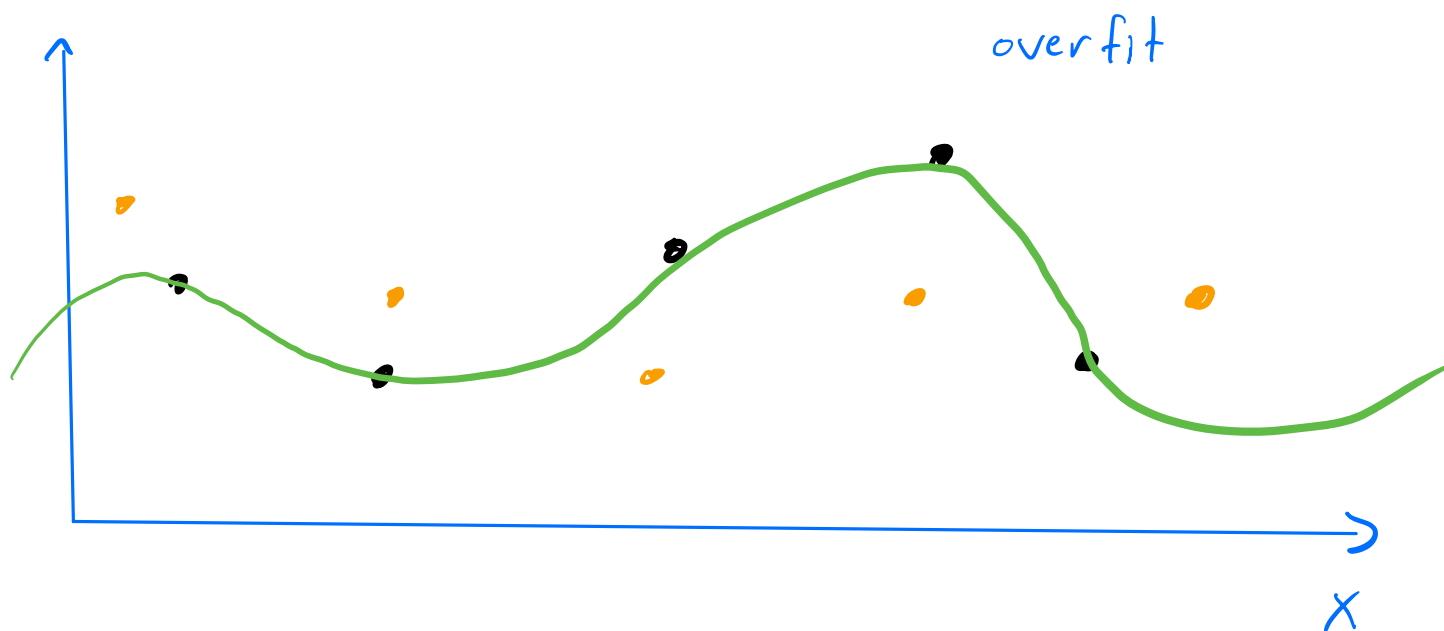
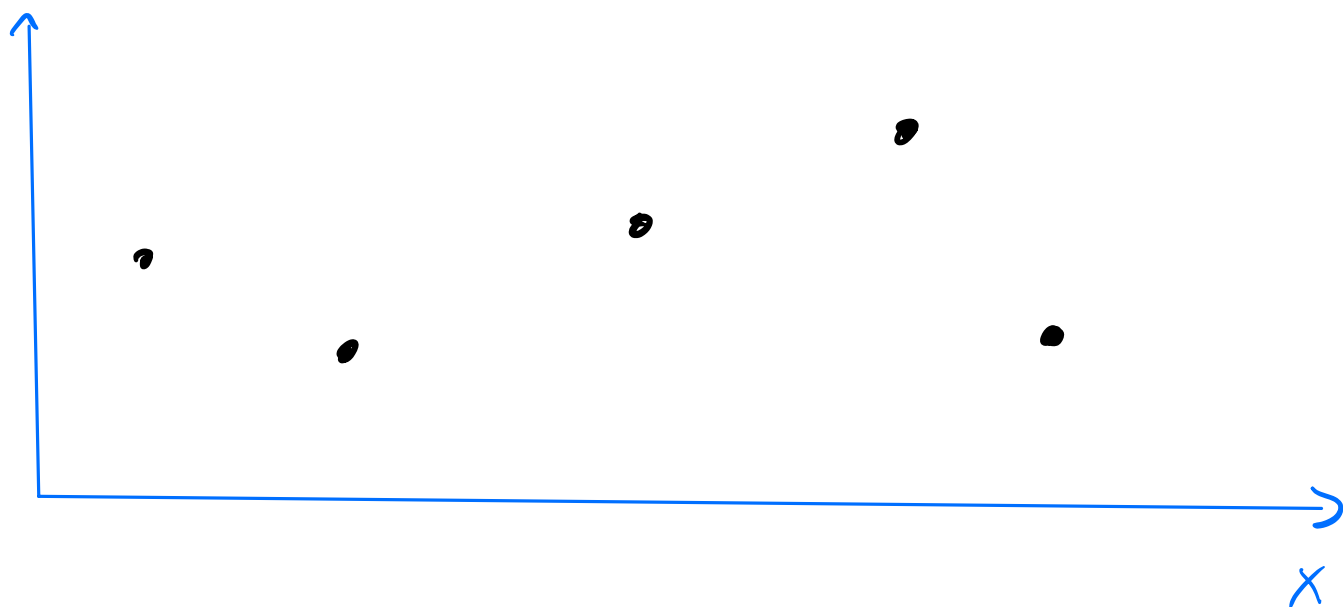
Pick $f \in \mathcal{F}$ then gather D
(i.e. f is chosen independently of D)



Then $E[\hat{L}(f)] = L(f)$

$$\text{Var}[\hat{L}(f)] = \frac{1}{n} \text{Var}[\ell(f(\vec{X}), Y)]$$

We are gathering D first and then
picking $\hat{f}_D \in \mathcal{F}$! (i.e. $\hat{f}_D$ depends on D)



In general

$$\mathbb{E}[\hat{L}(\hat{f}_0)] \neq \mathbb{E}[L(\hat{f}_0)]$$

$l(\hat{f}_0(\vec{X}_i), Y_i)$ are not i.i.d.

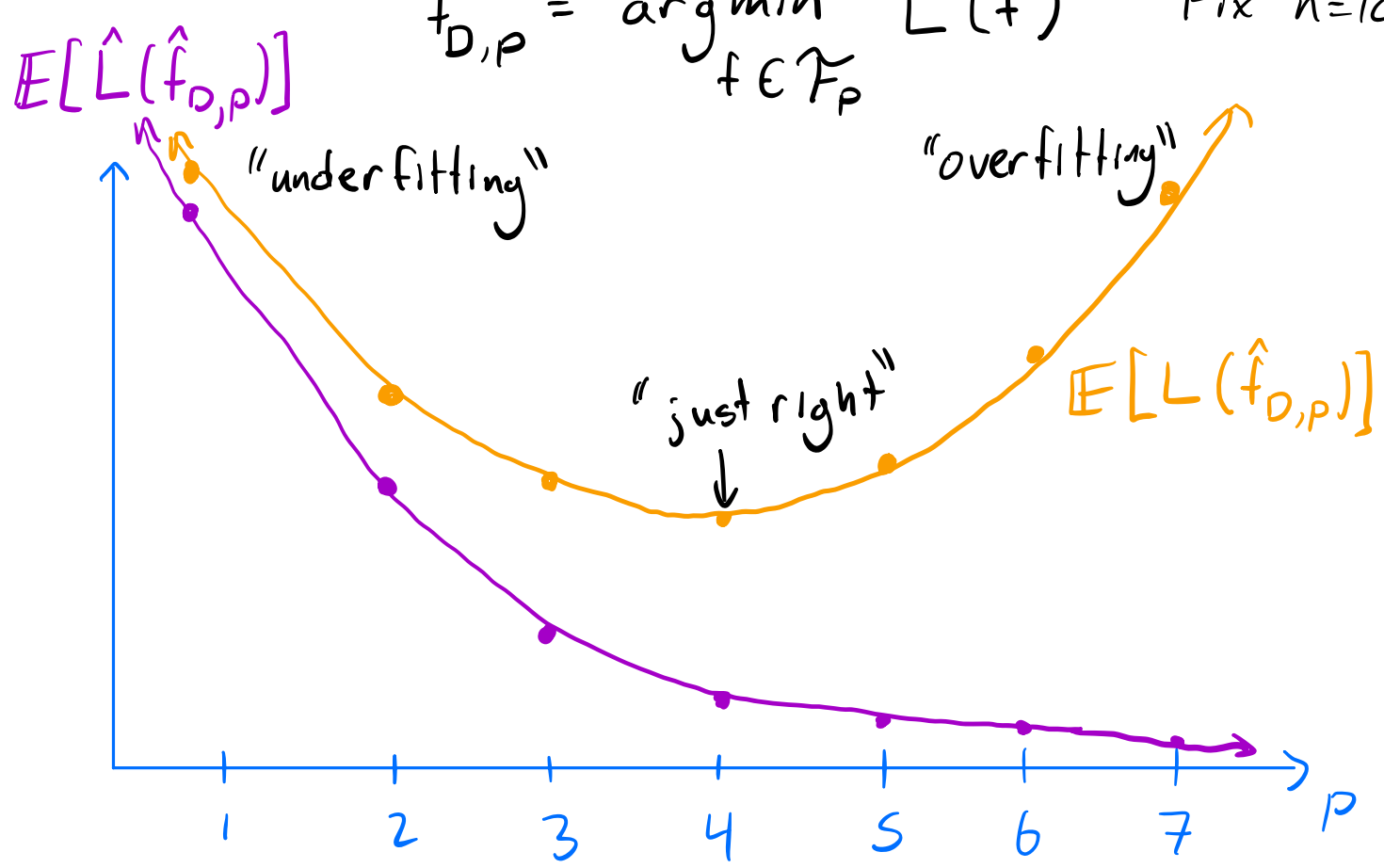
$\hat{f}_0$ depends $(\vec{X}_1, Y_1), \dots, (\vec{X}_n, Y_n)$

It can be shown that

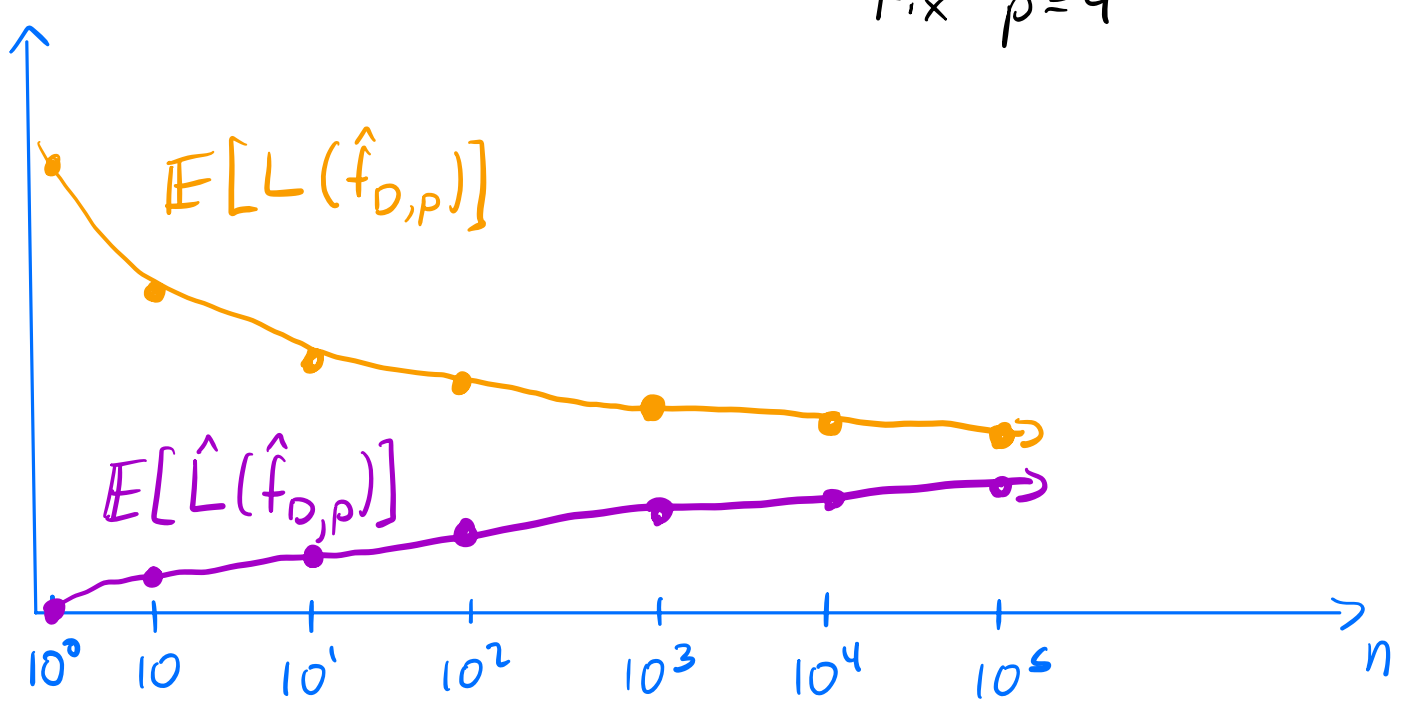
$$\mathbb{E}[L(\hat{f}_0)] - \mathbb{E}[\hat{L}(\hat{f}_0)] \text{ is:}$$

- non-negative (i.e. $\mathbb{E}[\hat{L}(\hat{f}_0)] \leq \mathbb{E}[L(\hat{f}_0)]$)
- increases as $\tilde{F}$ gets more complex
- decreases as n (# of data points) increases

$$\hat{f}_{D,p} = \operatorname{argmin}_{f \in \mathcal{F}_p} \hat{L}(f) \quad \text{Fix } n=100$$



Fix $p=4$



In practice we only have a fixed dataset $\mathcal{D}$

Train on a $\mathcal{D}_{\text{train}}$ and evaluate on $\mathcal{D}_{\text{test}}$:

$$\mathcal{D}_{\text{train}} = ((\tilde{X}_1, Y_1), \dots, (\tilde{X}_{n-m}, Y_{n-m}))$$

$$\mathcal{D}_{\text{test}} = ((\tilde{X}_{n-m+1}, Y_{n-m+1}), \dots, (\tilde{X}_n, Y_n))$$

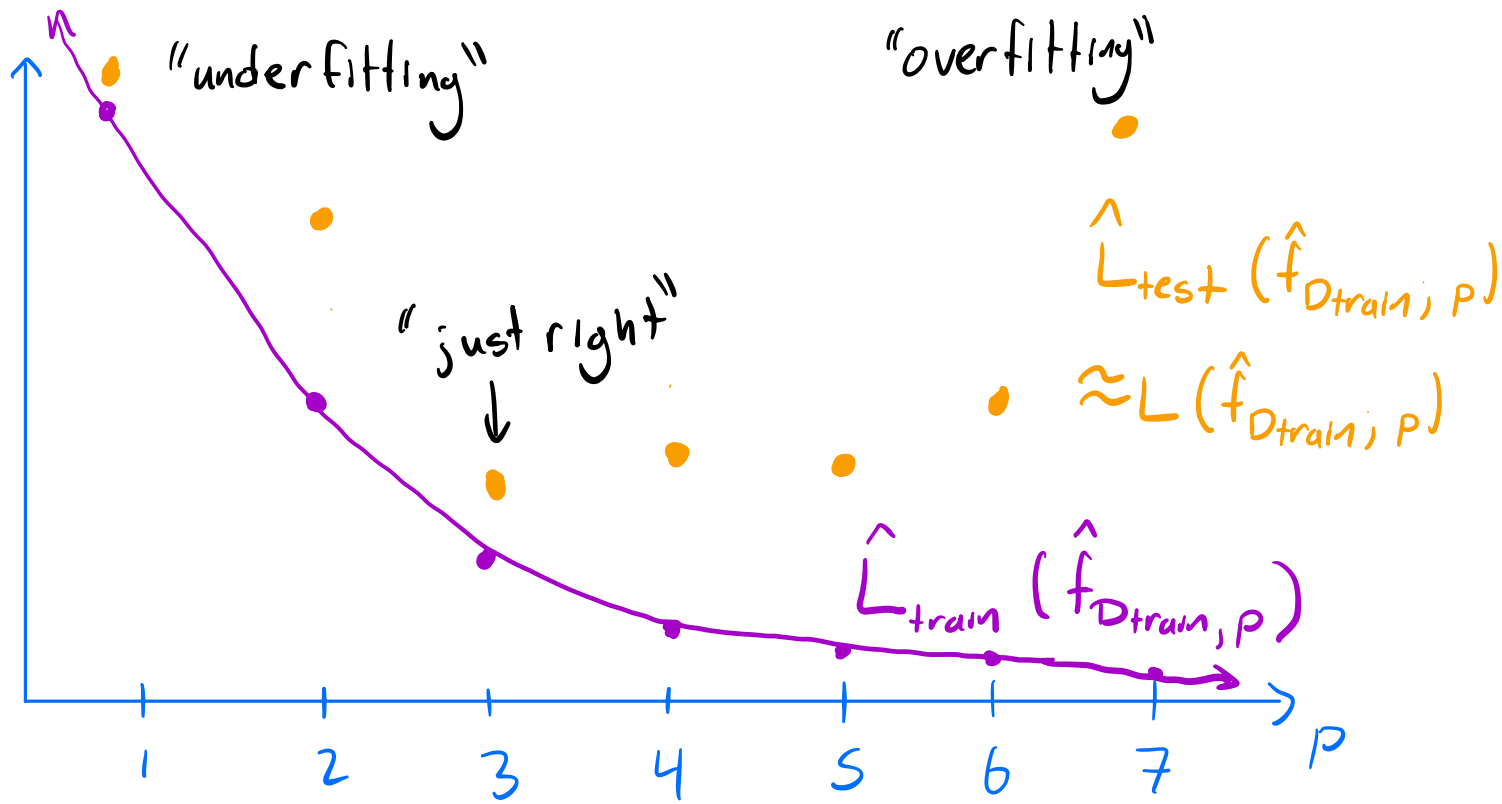
$$|\mathcal{D}_{\text{train}}| = n - m, |\mathcal{D}_{\text{test}}| = m$$

$$\text{Then } \mathcal{A}(\mathcal{D}_{\text{train}}) = \hat{f}_{\mathcal{D}_{\text{train}}} = \arg \min_{f \in \mathcal{F}} \hat{L}_{\text{train}}(f)$$

$$\hat{L}_{\text{train}}(\hat{f}_{\mathcal{D}_{\text{train}}}) = \frac{1}{n-m} \sum_{i=1}^{n-m} \ell(\hat{f}_{\mathcal{D}_{\text{train}}}(\tilde{X}_i), Y_i)$$

$$\hat{L}_{\text{test}}(\hat{f}_{\mathcal{D}_{\text{train}}}) = \frac{1}{m} \sum_{i=n-m+1}^n \ell(\hat{f}_{\mathcal{D}_{\text{train}}}(\tilde{X}_i), Y_i)$$

$$A_p(D_{\text{train}}) = \hat{f}_{D_{\text{train}}, p} = \arg \min_{f \in \mathcal{F}_p} \hat{L}_{\text{train}}(f)$$

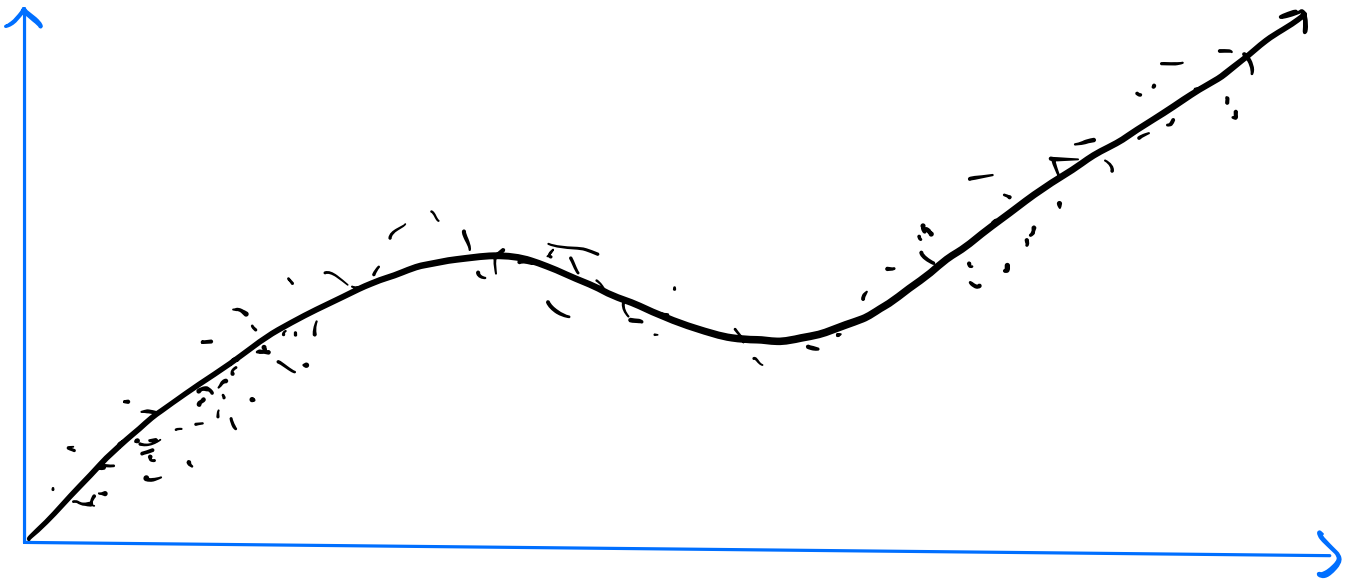


Decomposing $\mathbb{E}[L(A(D))]$ into Error Terms

• You know $P_{\bar{X}, Y}$

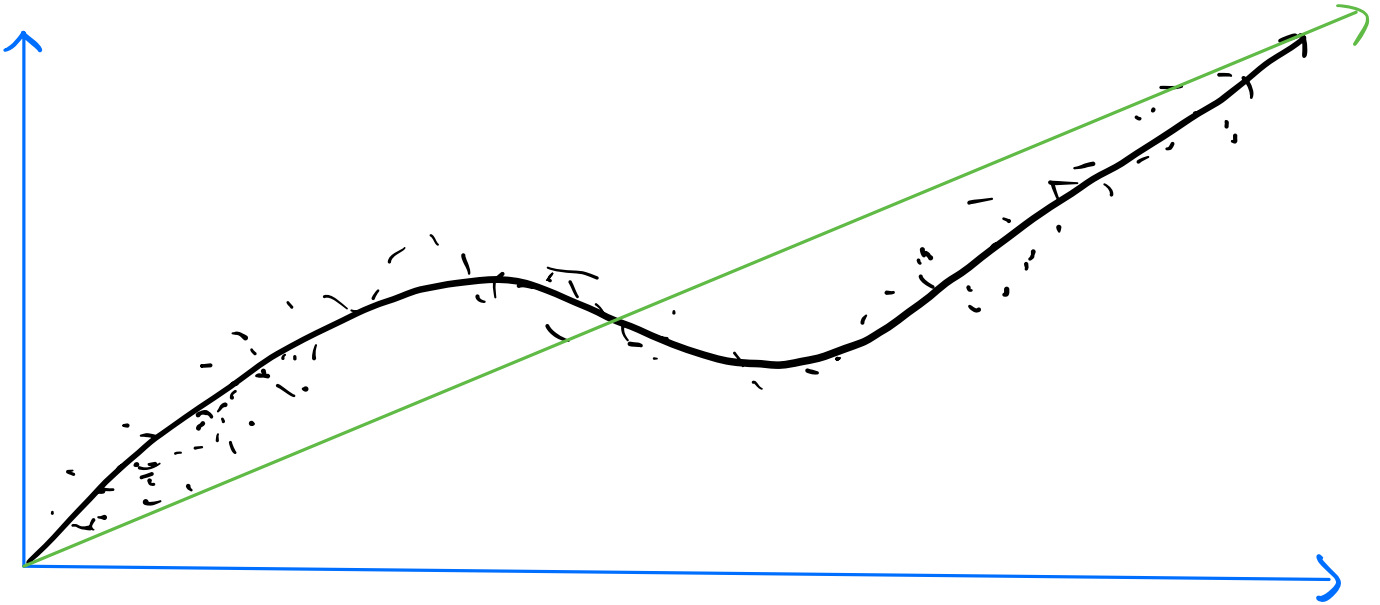
$$f_{\text{Bayes}} = \underset{f \in \{f | f: X \rightarrow Y\}}{\operatorname{argmin}} L(f)$$

Bayes optimal predictor



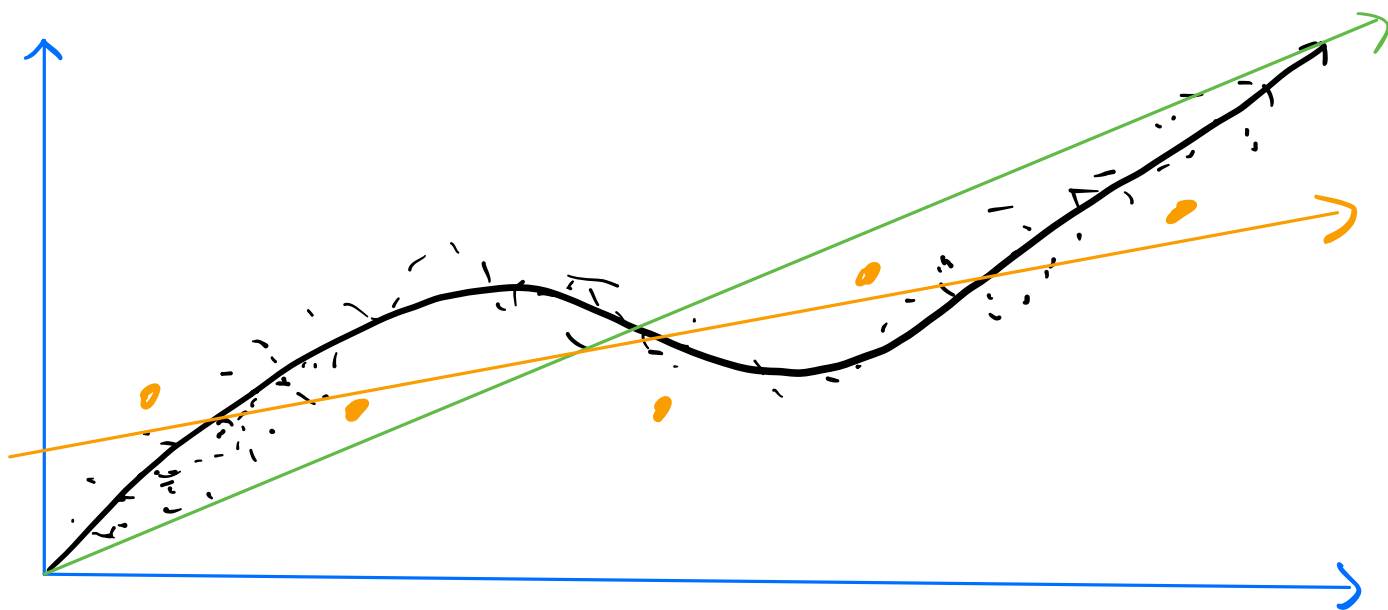
• You know $P_{\bar{X}, Y}$ but we have to output from $\mathcal{F}_1$

$$f^* = \underset{f \in \mathcal{F}_1}{\operatorname{argmin}} L(f)$$

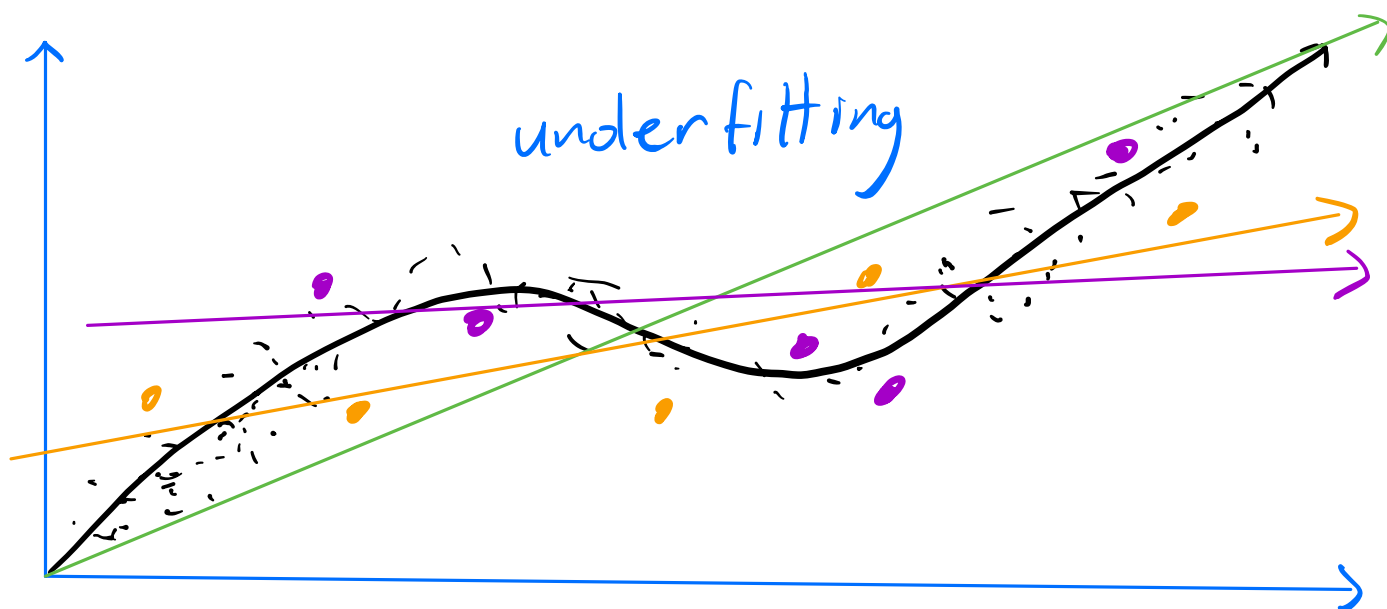


- You don't know $P_{\bar{x}, y}$ but we have to output from $\mathcal{F}_1$ and we have a dataset $\mathcal{D}_1$

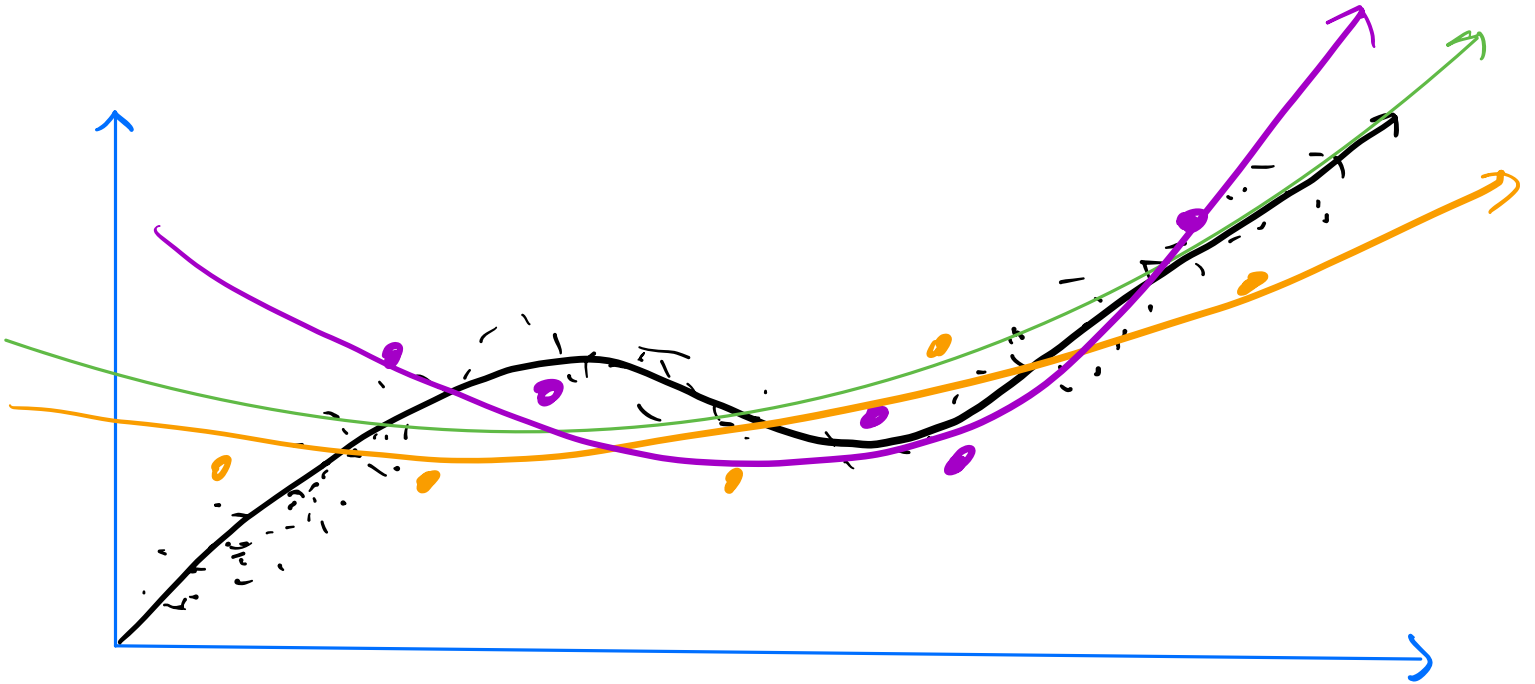
$$\hat{f} = \operatorname{argmin}_{f \in \mathcal{F}_1} \hat{L}(f)$$



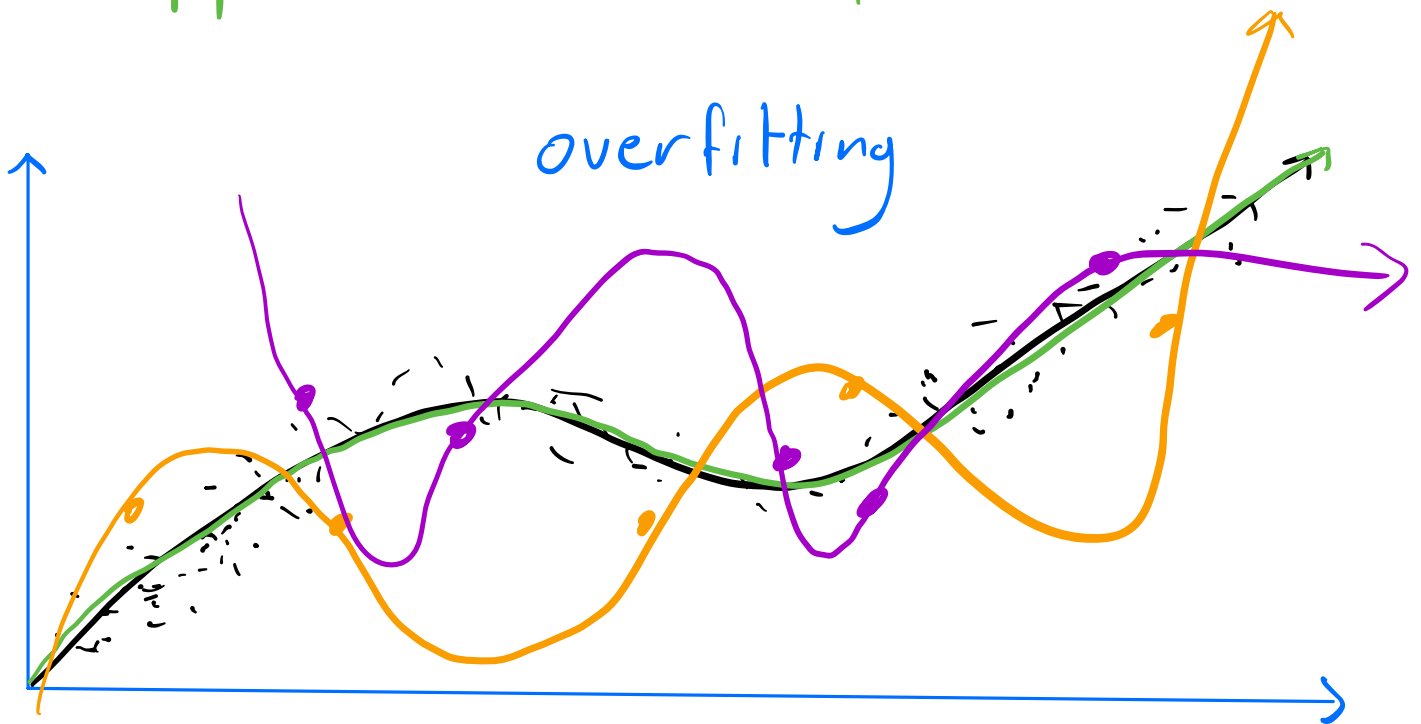
$\mathcal{D}_2$

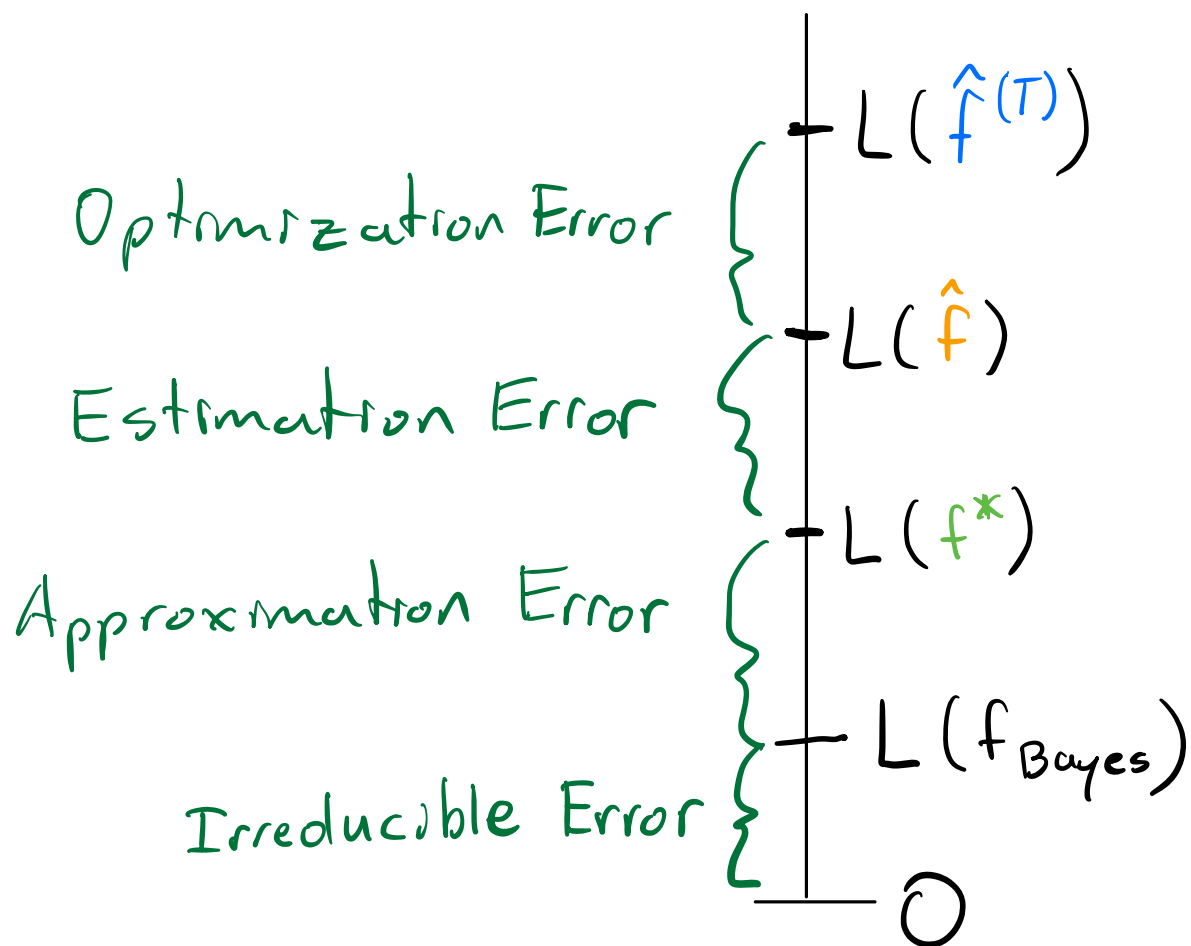


Suppose $\mathcal{F}_2$



Suppose $\mathcal{F}_{10}$ and $f_{\text{Bayes}} \in \mathcal{F}_{10}$





$$f_{\text{Bayes}} = \underset{f \in \{f | f: X \rightarrow Y\}}{\operatorname{argmin}} L(f) \quad f^* = \underset{f \in \mathcal{F}_i}{\operatorname{argmin}} L(f) \quad \hat{f}_D = \underset{f \in \mathcal{F}_i}{\operatorname{argmin}} \hat{L}(f)$$

$$\begin{aligned} E[L(\hat{f}_D)] &= \underbrace{E[L(\hat{f}_D)] - L(f^*)}_{\text{Estimation error (EE)}} \\ &\quad + \underbrace{L(f^*) - L(f_{\text{Bayes}})}_{\text{Approximation Error (AE)}} + \underbrace{L(f_{\text{Bayes}})}_{\text{Irreducible Error}} \end{aligned}$$

What affects the errors:

Irreducible Error: Due to inherent noise in the labels

- Can decrease if you gather more informative features

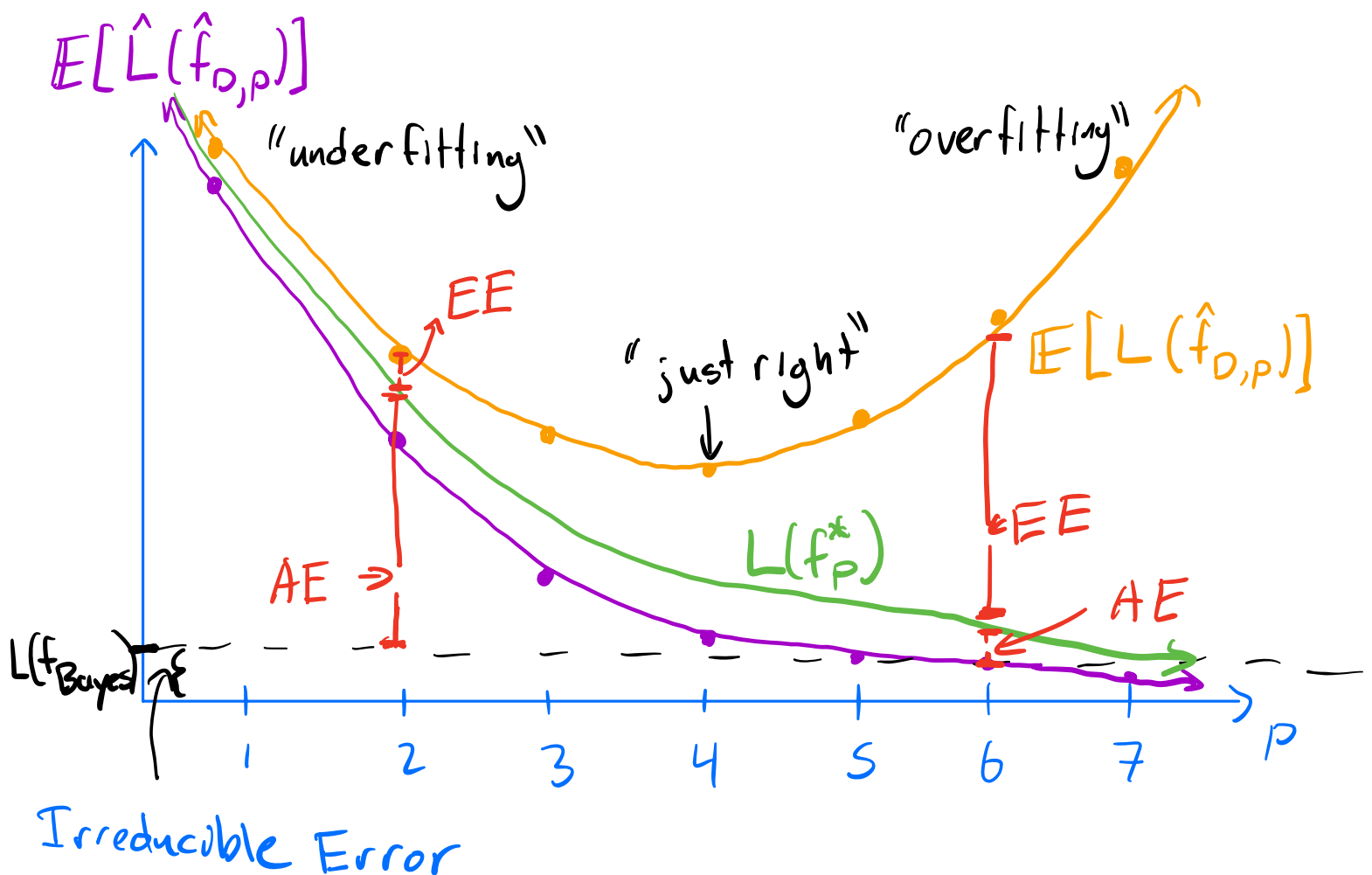
Approximation Error: Due to function class $\mathcal{F}$

- Decrease or stays the same if you make $\mathcal{F}$ larger

Estimation Error: Due to using a dataset D

- Decrease or stays the same if n increases
- Increase or stay the same if $\mathcal{F}$ gets larger

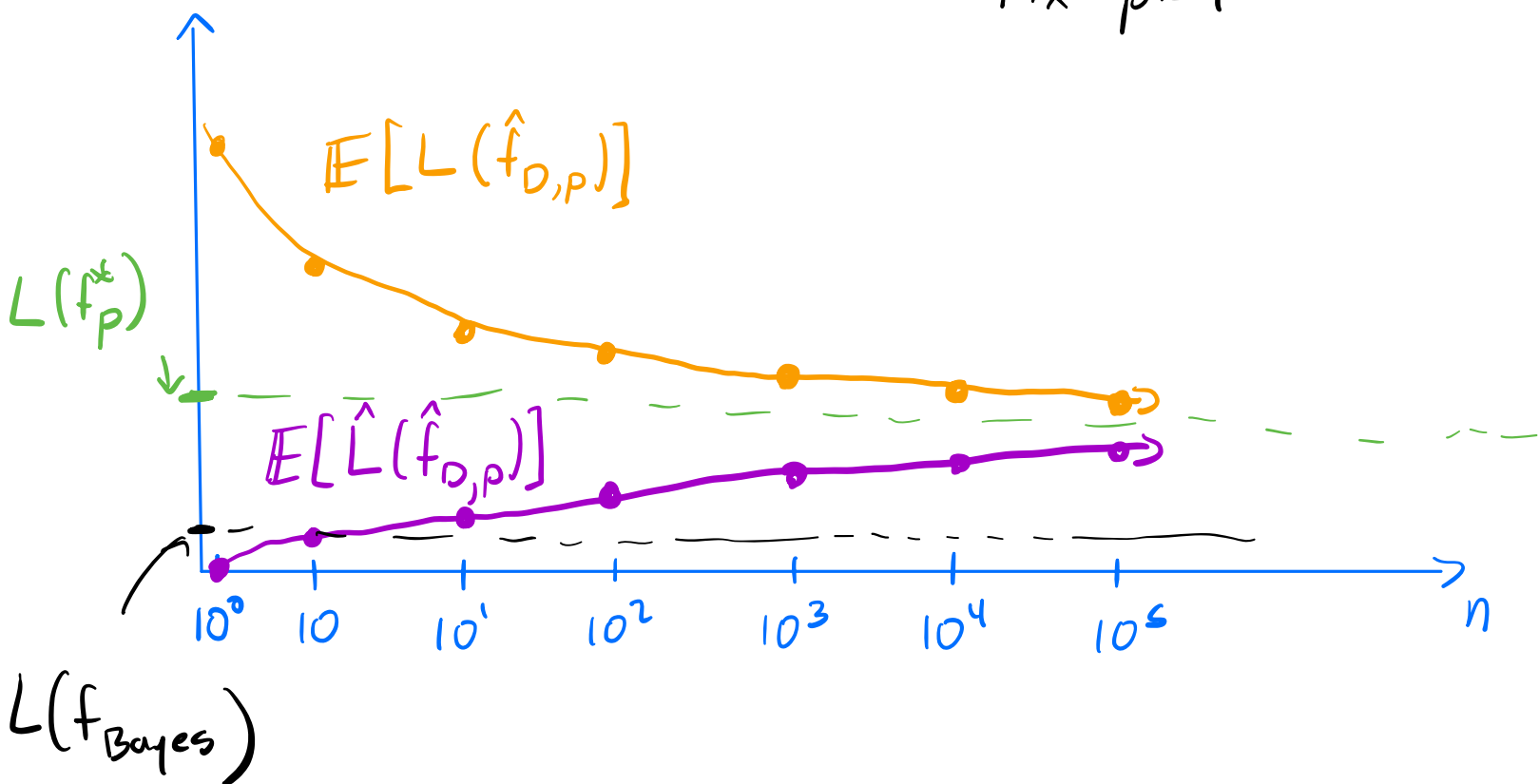
$$\hat{f}_{D,p} = \operatorname{argmin}_{f \in \mathcal{F}_p} \hat{L}(f) \quad \text{Fix } n=100$$



Underfitting: High AE, Low EE

Overfitting: Low AE, High EE

Fix $p=4$



Bias-Variance Decomposition

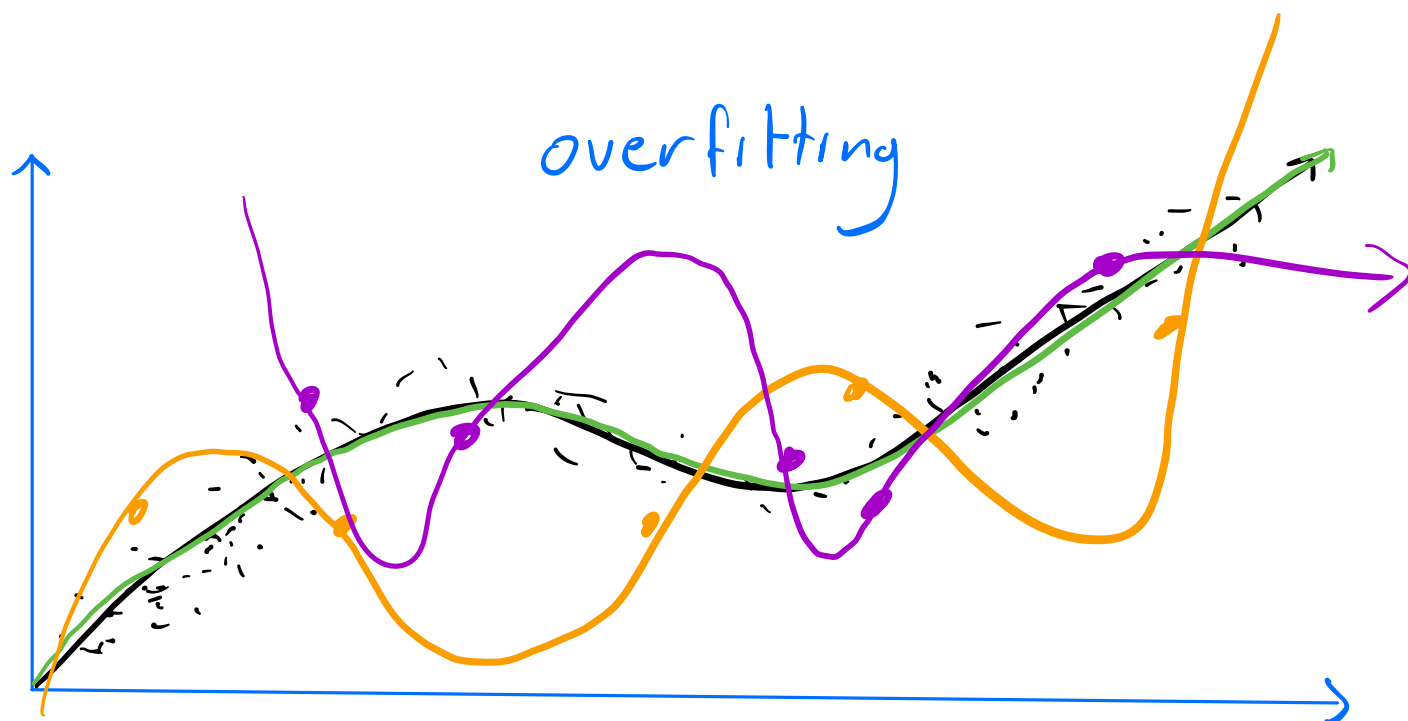
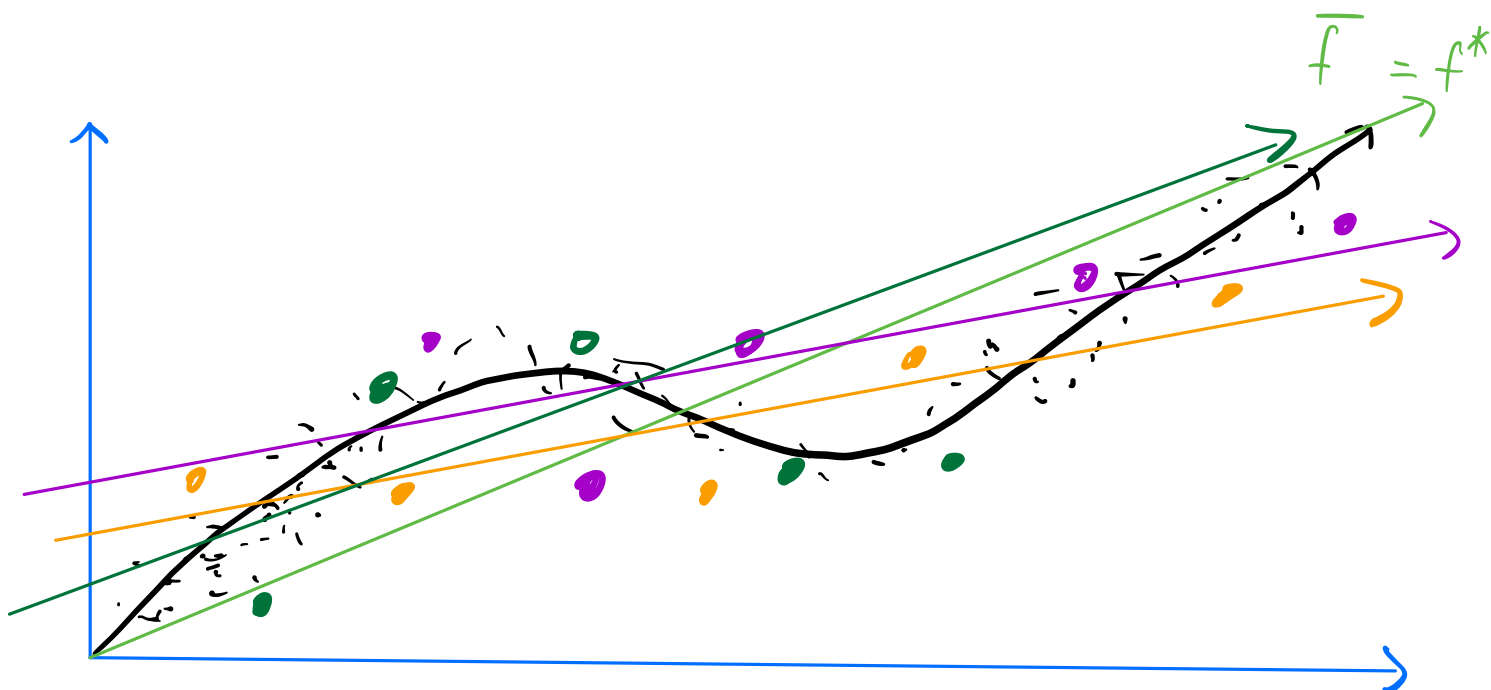
$$\mathbb{E}[L(\hat{f}_D)] = \mathbb{E}[L(\hat{f}_D)] - L(f_{\text{Bayes}}) + L(f_{\text{Bayes}})$$

regression and squared loss

$$= \underbrace{\mathbb{E}\left[\mathbb{E}\left[(\hat{f}_D(x) - \bar{f}(x))^2 \mid X\right]\right]}_{\text{Variance} = \text{EE}} + \underbrace{\mathbb{E}\left[(\bar{f}(X) - f_{\text{Bayes}}(X))^2\right]}_{\text{Bias} = \text{AE}} + L(f_{\text{Bayes}})$$

$$\bar{f}(x) = \mathbb{E}[\hat{f}_D(x) \mid X=x] \text{ expected predictor}$$
$$= f^*(x)$$

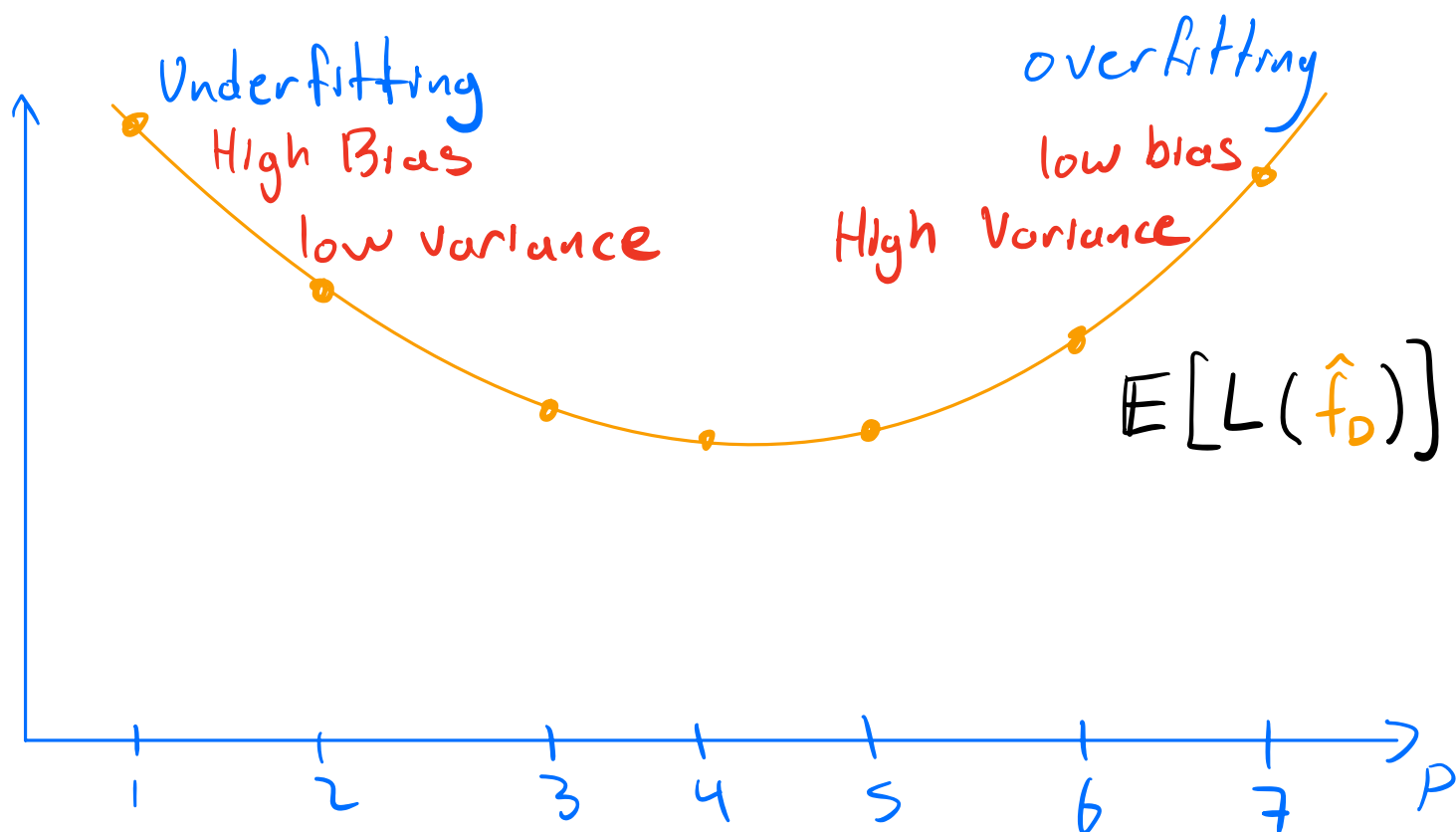
$$\hat{f} = \underset{f \in \mathcal{F}_1}{\operatorname{argmin}} \hat{L}(f)$$



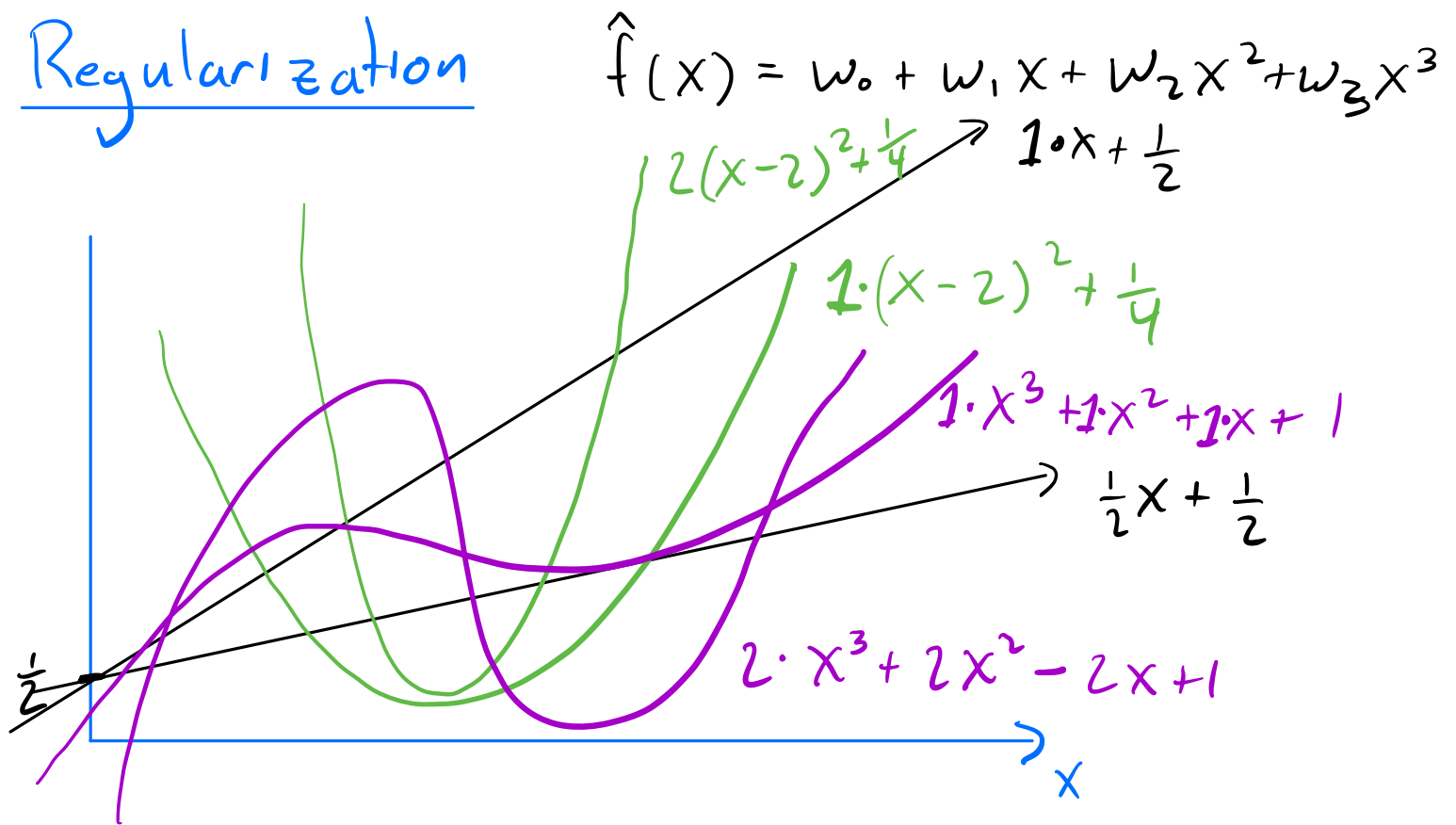
Bias $\downarrow$ if $\mathcal{F} \uparrow$

Variance $\downarrow$ if $n \uparrow$

$\uparrow$ if $\mathcal{F} \uparrow$



Regularization



let $\vec{w} = (w_0, \dots, w_{\bar{p}-1})^T \in \mathbb{R}^{\bar{p}}$

Observation: large values of $w_0, \dots, w_{\bar{p}-1}$ leads to a more complex $f_p(\vec{x}) = \phi_p(\vec{x})^T \vec{w}$

$$\widetilde{F}_{p,c} = \left\{ f \mid f: \mathbb{R}^{d+1} \rightarrow \mathbb{R}, \text{ s.t. } f(\vec{x}) = \phi_p(\vec{x})^T \vec{w}, \right. \\ \left. \vec{w} \in \mathbb{R}^{\bar{p}}, \text{ and } \sum_{j=1}^{\bar{p}-1} w_j^2 \leq C, C > 0 \right\}$$

$$A_{p,\lambda}(D) = \hat{f}_{p,\lambda} \text{ where } \hat{f}_{p,\lambda}(\bar{x}) = \phi_p(\bar{x})^T \hat{\vec{w}}_\lambda$$

regularizer

$$\hat{\vec{w}}_\lambda = \argmin_{\vec{w} \in \mathbb{R}^{\bar{p}}} \left(\underbrace{\frac{1}{n} \sum_{i=1}^n \ell(\phi_p(\vec{x}_i)^T \vec{w}, y_i)}_{g(\vec{w})} + \underbrace{\frac{\lambda}{n} \sum_{j=1}^{\bar{p}-1} w_j^2}_{\text{regularizer}} \right)$$

"Ridge Regression"

$g(\vec{w})$ convex!

doesn't
use
 w_0

$\lambda \in (0, \infty)$ regularization parameter

$$\hat{L}_\lambda(\vec{w}) = \underbrace{\frac{1}{n} \sum_{i=1}^n \ell(\phi_p(\vec{x}_i)^T \vec{w}, y_i)}_{\hat{L}(\vec{w})} + \frac{\lambda}{n} \sum_{j=1}^{\bar{p}-1} w_j^2$$

set $p=1 \Rightarrow \bar{p}=d+1$

$$\hat{\vec{w}}_\lambda = \argmin_{\vec{w} \in \mathbb{R}} \hat{L}_\lambda(\vec{w})$$

$$\vec{w}^{(t+1)} = \vec{w}^{(t)} - \eta^{(t)} \nabla \hat{L}_\lambda(\vec{w}^{(t)})$$

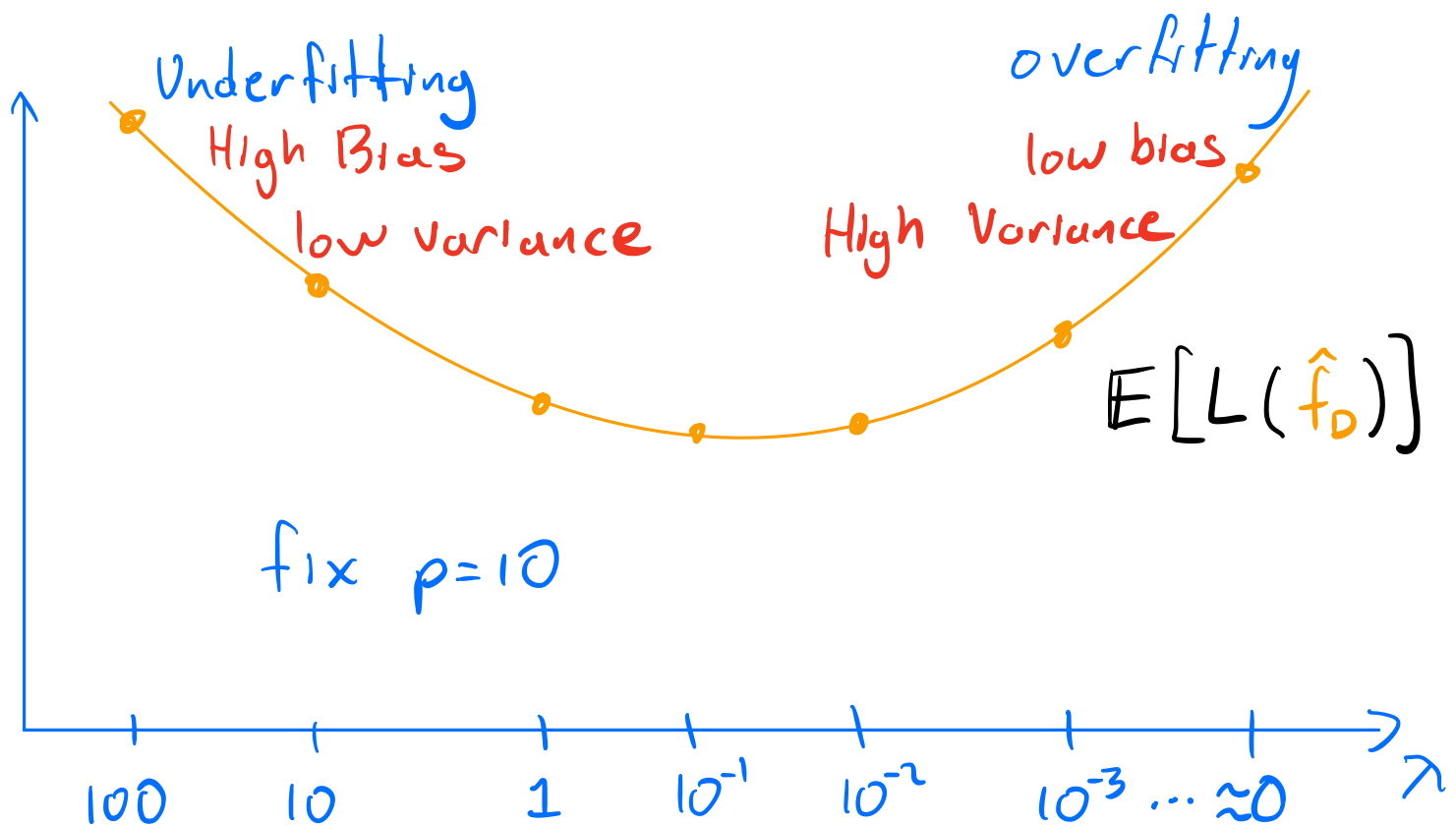
$$\nabla \hat{L}_\lambda = \left(\frac{\partial \hat{L}_\lambda}{\partial w_0}, \frac{\partial \hat{L}_\lambda}{\partial w_1}, \dots, \frac{\partial \hat{L}_\lambda}{\partial w_d} \right)^T$$

$$\frac{\partial \hat{L}_\lambda}{\partial w_0} = \frac{\partial \hat{L}}{\partial w_0} = \frac{2}{n} \sum_{i=1}^n (\vec{x}_i^T \vec{w} - y_i) x_{i0}$$

for $k \in \{1, \dots, d\}$

$$\frac{\partial \hat{L}_\lambda}{\partial w_k} = \frac{\partial \hat{L}}{\partial w_k} + \frac{2 \left(\frac{\lambda}{n} \sum_{j=1}^d w_j^2 \right)}{2 w_k}$$

$$= \frac{2}{n} \sum_{i=1}^n (\vec{x}_i^T \vec{w} - y_i) x_{ij} + \frac{2\lambda}{n} w_k$$



Algorithm: BGD Linear Regression Learner
(with a constant step size, and regularization)

input: $D = ((\vec{x}_1, y_1), \dots, (\vec{x}_n, y_n))$, η , T , λ

$\vec{w}$ = random vector in $\mathbb{R}^{d+1}$

for $t = 1, \dots, T$

$$\nabla \hat{L}(\vec{w}) = \frac{2}{n} \sum_{i=1}^n (\vec{x}_i^T \vec{w} - y_i) \vec{x}_i$$

$$\nabla \hat{L}(\vec{w})[1:d] += \frac{2\lambda}{n} \vec{w}[1:d]$$

$$\vec{w} = \vec{w} - \eta \nabla \hat{L}(\vec{w})$$

return $\hat{f}: \mathbb{R}^{d+1} \rightarrow \mathbb{R}$ where $\hat{f}(\vec{x}) = \vec{x}^T \vec{w}$

if $\lambda \uparrow$ Bias $\uparrow$

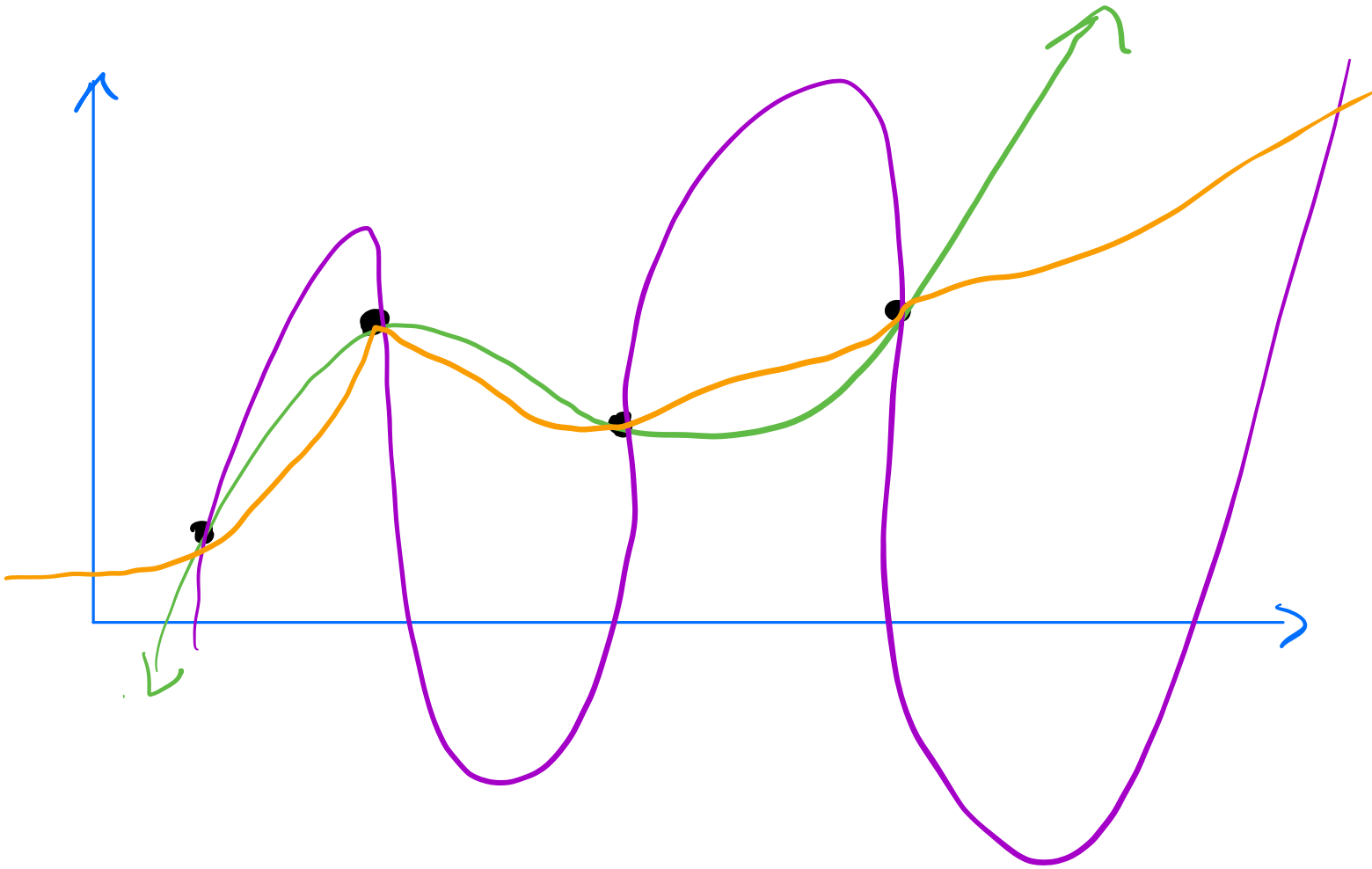
Var $\downarrow$

AE is same

EE is complicated

if $\lambda \neq 0$ then $\bar{f} \neq f^*$

Double Descent (Not graded)



$$\hat{f} \in \arg \min_{f \in \mathcal{F}} \hat{L}(f)$$

$n = 8$

