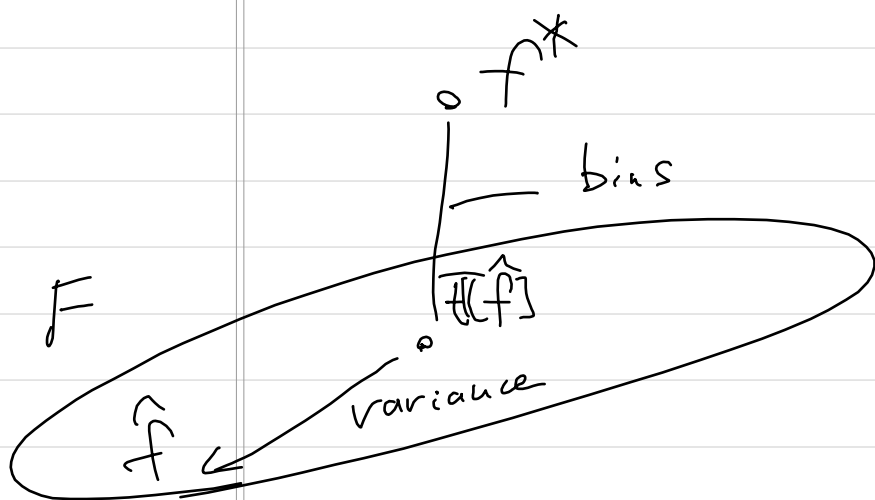


(a) Under-fitting : condition that statistical model cannot capture feature and rules

Over-fitting : condition that statistical model fits to not only feature but also noise

Bias : gap between expectation of estimate and true model



$R(f)$ is risk function

$$f^* = \underset{f}{\operatorname{argmin}} R(f)$$

$$\hat{f} = \underset{f \in F}{\operatorname{argmin}} R(f)$$

$$\text{Bias}[\hat{f}] = \mathbb{E}[\hat{f}] - f^*$$

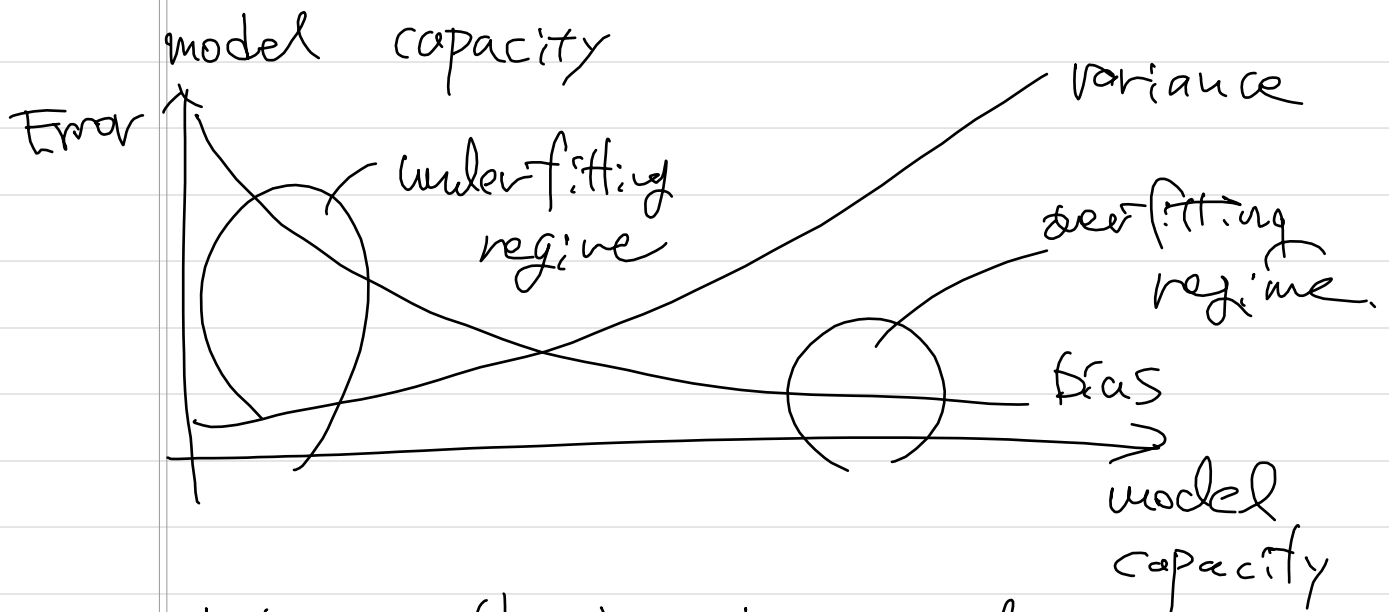
$$\text{Variance}[\hat{f}] = \mathbb{E}[\hat{f}^2] - (\mathbb{E}[\hat{f}])^2$$

20A_Q1

2

Bias: an error from erroneous assumption in the learning algorithm, high bias can cause an algorithm to miss the relevant relations between features and target outputs (underfitting)

Variance: an error from sensitivity to small fluctuations in the training set. High variance may result from an algorithm modeling the random noise in the training data (overfitting)



that control the bias - and - variance trade-off,

21A-Q1

3

(cc) #Train sample $\nearrow$ variance $\searrow$?

[2] (a) Diff of param hyperparam

updated by training data set

chose by evaluation on
validation data
as model selection

which determine training behaviour
and efficient capacity

$\downarrow$ param

$\downarrow$ hyperparam

(b) (1) Linear Regression

$$\theta = \{w + b\}$$

$$\text{obj} = l(f(x), y) = \frac{1}{2} \|f(x) - y\|^2 + \lambda \|w\|^2$$

$$f(x) = w^T x + b$$

$$w \in \mathbb{R}^{d \times d}, y \in \mathbb{R}^d, x \in \mathbb{R}^d, b \in \mathbb{R}^d.$$

(2)
 λ : regularization

loss func

$$\theta_{t+1} = \theta_t - \eta \nabla_{\theta} l(\hat{y}, y)$$

$\uparrow$
 η = step size

(2) MLP

$$\hat{y} = f(x; \theta)$$

$\theta \in \mathbb{R}^d$ model param

2(A-Q)

4

$\theta \sim \text{Beta}(\alpha, \beta) \leftarrow \text{hyperparam}$)
 $X_i \sim \text{Bern}(\theta) \leftarrow \text{param}$,

$$y \sim N(xb, \sigma^2 I)$$

$$\hat{b} = (X^T X + kI)^{-1} X^T y$$

$b, \theta \leftarrow \text{param}$
 $k \leftarrow \text{hyperparam}$

3

1. Predict likelihood

input quiz score
output likelihood.

- ① regression (logistic)
- ② MLP w/ sigmoid

2. Predict height

- ① linear regression
regression tree

3. Unsupervised classification (text)

non-hierarchical
 k-means (clustering)
 Hierarchical clustering (Ward method)

[4]

Cross Entropy and MLE (relationship)

if we want to maximize log likelihood

$$\hat{\theta} = \underset{\theta}{\operatorname{argmax}} p(y|x;\theta)$$

$$\Leftrightarrow \hat{\theta} = \underset{\theta}{\operatorname{argmax}} \log p(y|x;\theta)$$

$$= \underset{\theta}{\operatorname{argmax}} L(\theta)$$

$$= \underset{\theta}{\operatorname{argmax}} \sum_{k=1}^K y_k \log p_k$$

$$= \underset{\theta}{\operatorname{argmin}} \underbrace{- \sum_{k=1}^K y_k \log p_k}$$



it is equivalent $CE(y_k, p_k)$

we focus
 on train
 sample

$$\frac{1}{K} \sum_{k=1}^K p(y_k|x;\theta)^{y_k}$$

2(A-Q1)

6

$$\boxed{5} \quad \mathcal{N} = \{(x_1, y_1), \dots, (x_n, y_n)\} \sim \mathcal{D}$$

$$x \in \mathbb{R}^d \times \{0, 1\}$$

(a) $\boxed{A \mid B} \leftarrow \text{data set}$

step 1 use A as train data
then evaluate $\hat{\theta}_A$ on dataset B

step 2 use B as vice-versa

$$\text{estimated error} = \frac{1}{2} (l(B, \hat{\theta}_A) + l(A, \hat{\theta}_B))$$

(b) CV error of learned classifier

an biased estimate of its true error?