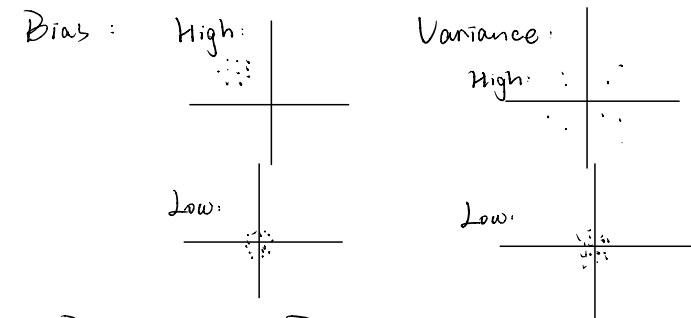


Bias and Variance

1. Generalization Error



Regression:

$$D: \{x_i, y_i\}_{i=1}^N$$

Determine the weight w .

$$\Rightarrow y = w^T \phi(x)$$

New Statement:

$$S = \{(x_i, y_i)\}_{i=1}^m, \text{ where } y_i \in \mathbb{R}$$

A hypothesis class $\mathcal{H}$ of regression functions.

e.g. $\mathcal{H} = \{h(x) = w^T \phi(x) | w\}$. "假设空间"

$\Rightarrow$ Choose one hypothesis h from $\mathcal{H}$.

Training Error

$$\hat{E}(h) = \frac{1}{m} \sum_{i=1}^m (y_i - h(x_i))^2$$

Training:

$$\text{Objective: } \hat{h} = \arg \min_{h \in \mathcal{H}} \hat{E}(h)$$

$\Rightarrow$ training error: $\hat{E}(\hat{h})$

Generalization Error:

$$E(h_s) = E_{(x,y) \sim D} [(y - h(x))^2]$$

Expectation (sampling (x,y) from D)

Generalization Gap:

$$E(h_s) - E(\hat{h}_s)$$

$$E = E_s[E(h_s)] = E_s[E_{(x,y) \sim D} [(y - h(x))^2]]$$

E : Expected generalization error $\rightarrow$ We want to minimize it!

★ The Bias-Variance Decomposition

$$E = E_s[E_{(x,y)} [(y - h_s(x))^2]]$$

1. $h_s(x) \rightarrow h_s$ for readability

$$E = E_s[E_{(x,y) \sim D} [(y - h_s)^2]]$$

$$= E_s E_{(x,y)} [(y - E_s(h_s) + E_s(h_s) - h_s)^2]$$

$$= E_s E_{(x,y)} [(y - E_s(h_s))^2] + E_s E_{(x,y)} [(E_s(h_s) - h_s)^2]$$

$$+ 2 E_s E_{(x,y)} [(y - E_s(h_s)) (E_s(h_s) - h_s)]$$

Scalar (for E_s)

$$= E_s E_{(x,y)} [(y - E_s(h_s))^2] + E_s E_{(x,y)} [(E_s(h_s) - h_s)^2]$$

$$+ 2 E_{(x,y)} [(y - E_s(h_s)) \cdot (E_s(E_s(h_s)) - E_s(h_s))]$$

$$\star E_s(E_s(h_s)) = E_s(h_s)!$$

$$\Rightarrow E_s(E_s(h_s)) - E_s(h_s) = 0!$$

$$= \frac{E_s E_{(x,y)} [(y - E_s(h_s))^2]}{(\text{Bias})^2} + E_s E_{(x,y)} [(E_s(h_s) - h_s)^2]$$

$$= \frac{E_{(x,y)} [(y - E_s(h_s))^2]}{(\text{Bias})^2} + E_x E_s [(E_s(h_s) - h_s)^2]$$

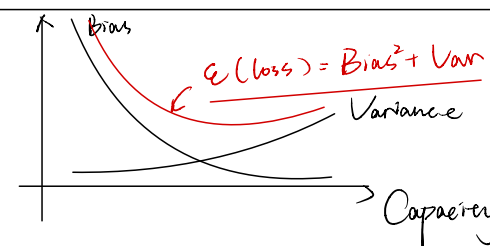
"(Bias)²" $\leftarrow$ "Trade-off" $\rightarrow$ "Variance"
"偏差程度" $\leftarrow$ "离散程度"

$$\star \text{Expected Generalization Error} = \text{Bias}^2 + \text{Variance}!$$

illustration:

Impact of Model Capacity on Bias & Variance

$$\star \text{Model capacity } \uparrow, \text{ Variance } \uparrow \& \text{ Bias } \downarrow$$



Impact of Regularization on Bias & Variance

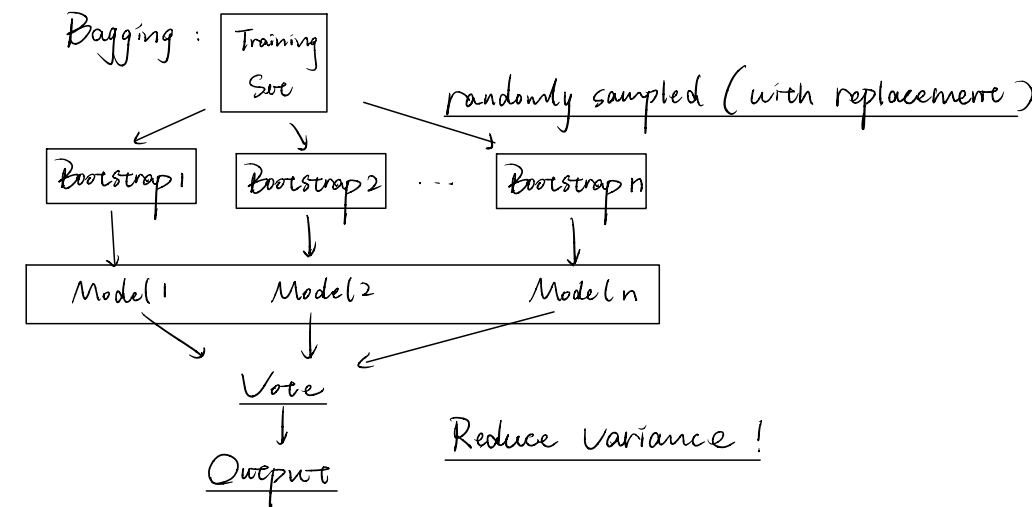
Since: Regularization $\uparrow$ Model Capacity $\downarrow$

$$\star \text{Regularization } \uparrow, \text{ Variance } \downarrow \& \text{ Bias } \uparrow$$

Do not over-regularize

Ensemble Learning

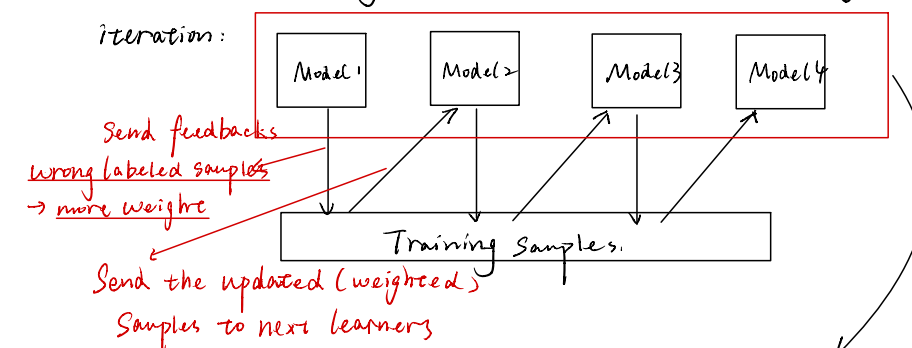
Bagging: reduce variance
Boosting: reduce bias.



Boosting:

init: All training sample assigned same weight

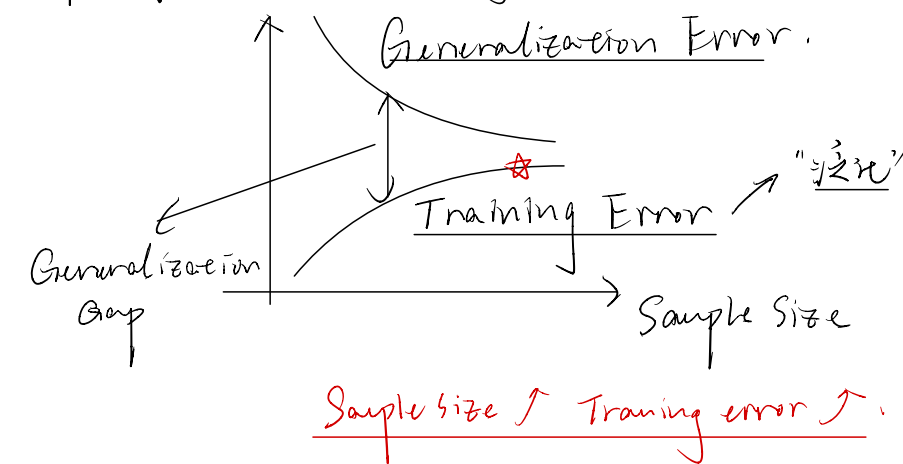
iteration:



final outcomes: final model = weighted sum of all learners.

"错误率低的模型权重更大"

Impact of Sample Size on Training/Generalization Error:



Impact of Sample Size on Bias & Variance:

样本量	对方差(Variance)的影响	对偏差(Bias)的影响	主要风险
非常小	极高 (模型不稳定, 随数据波动大)	较低 (取决于模型, 但通常能拟合训练点)	过拟合
中等	中等 (开始收敛)	开始反映真实模型能力	取决于模型复杂度
非常大	极低 (模型稳定, 泛化能力强)	收敛到固定值 (若模型太简单则偏差高)	欠拟合 (如果存在)

Scaling Law:

Performance of LLM increase with:

1. Sample Size
2. Model Capacity
3. GPU Power