

Model Comparison + Information Criteria

⇒ Mallows' C_p

bias/variance tradeoff

prediction error

Prediction Problem: Given data $\underline{y}$, $X = \begin{pmatrix} x_1 & \dots & x_p \\ 1 & & 1 \end{pmatrix}$

- 1) Obtain OLS estimate $\hat{\beta}$
- 2) Construct fitted values $\hat{\underline{u}} = X\hat{\beta} \in \mathbb{R}^{n \times 1}$
- 3) Use $\hat{\underline{u}}$ to predict $\underline{y}^*$, a new data vector generated under the same conditions as $\underline{y}$, the original data vector.

this means $E[\underline{y}] = E[\underline{y}^*] = \underline{\mu}$, (might not be $X\beta$!)

Q: How well does $\hat{\underline{u}}$ do at predicting $\underline{y}^*$?

A: Evaluate with expected prediction sum of squares

$$PSS = \|\underline{y}^* - \hat{\underline{u}}\|^2, \text{ evaluate } \underline{E[PSS]}$$

Setup: $\underline{y} = \underline{\mu} + \underline{\epsilon}$ $V(\underline{\epsilon}) = V(\underline{\epsilon}^*) = \sigma^2 \mathbf{I}$, $\underline{\epsilon}, \underline{\epsilon}^*$ uncorrelated.
 $\underline{y}^* = \underline{\mu} + \underline{\epsilon}^*$

Given X , $\hat{\beta} = (X^T X)^{-1} X^T \underline{y}$
 $\hat{\underline{u}} = X\hat{\beta} = X(X^T X)^{-1} X^T \underline{y} \equiv P \underline{y}$ (P is a projection matrix)

IDENTITY #1 : $E[\|\hat{y} - \hat{\mu}\|^2] = n\sigma^2 + p\sigma^2 + \|(\mathbf{I} - \mathbf{P})\mu\|^2$

$\uparrow$ sampling variability $\uparrow$ estimation var $\uparrow$ bias²

Proof :

$$\begin{aligned} \|\hat{y} - \hat{\mu}\|^2 &= \|\mu + \epsilon^* - \hat{\mu}\|^2 \\ &= \|\epsilon^* + (\mu - \hat{\mu})\|^2 \\ &= \|\epsilon^*\|^2 + \|\mu - \hat{\mu}\|^2 + \epsilon^{*T}(\mu - \hat{\mu}) \end{aligned}$$

$$\begin{aligned} E[\text{"}] &= E[\|\epsilon^*\|^2] + E[\|\mu - \hat{\mu}\|^2] + E[\epsilon^{*T}(\mu - \hat{\mu})] \\ &= n\sigma^2 + E[\|\mu - \hat{\mu}\|^2] + 0 \\ &= n\sigma^2 + E[\|\mu - \hat{\mu}\|^2] \end{aligned}$$

$$\begin{aligned} \|\mu - \hat{\mu}\|^2 &= \|\mu - \mathbf{P}_y\|^2 = \|\mu - \mathbf{P}\mu - \mathbf{P}\epsilon\|^2 \\ &= \|(\mathbf{I} - \mathbf{P})\mu - \mathbf{P}\epsilon\|^2 \\ &= \|(\mathbf{I} - \mathbf{P})\mu\|^2 + \|\mathbf{P}\epsilon\|^2 - 2\epsilon^T \mathbf{P}(\mathbf{I} - \mathbf{P})\mu \\ &= \|(\mathbf{I} - \mathbf{P})\mu\|^2 + \epsilon^T \mathbf{P} \epsilon - 0 \\ E(\text{"}) &= \|(\mathbf{I} - \mathbf{P})\mu\|^2 + p\sigma^2 \end{aligned}$$

So $E[\|\hat{y} - \hat{\mu}\|^2] = n\sigma^2 + p\sigma^2 + \|(\mathbf{I} - \mathbf{P})\mu\|^2$

- ① All else being equal, prediction error is necessary in p
- ② Suppose $\mu = X\beta$. Then $(\mathbf{I} - \mathbf{P}_X)\mu = X\beta - X(X^T X)^{-1}X^T X\beta = \underline{0}$ Bias is zero.
- ③ Suppose $X = [X_1, X_2]$ but $\hat{\mu} = X_1\beta$ (too many predictors)
 $X_1 \in \mathbb{R}^{n \times p_1}, X_2 \in \mathbb{R}^{n \times p_2}$

if we use $\hat{u} = X_1 (X_1^T X_1)^{-1} X_1^T y$, then
 $= P_1 y$

$$E(\|y^* - \hat{u}\|^2) = n\sigma^2 + p_1\sigma^2 + \|\cancel{(I - P_1)}^0 u\|^2 \\ = \underline{n\sigma^2 + p_1\sigma^2}$$

if we use $\hat{u} = X(X^T X)^{-1} X^T y$ ($X = (X_1, X_2)$, too many predictors)
 $= P_X y$

$$E(\|y^* - \hat{u}\|^2) = n\sigma^2 + (p_1 + p_2)\sigma^2 + \|\cancel{(I - P_X)}^0 u\|^2 \\ = \underline{n\sigma^2 + (p_1 + p_2)\sigma^2}$$

So adding unnecessary predictors increases the prediction error.

④ If $u = [X \ X^*] \beta = \tilde{X} \beta$ (not enough predictors)

$$\text{then } E(\|y^* - \hat{u}\|^2) = n\sigma^2 + np^2 + \|\cancel{(I - P_X)}^{\uparrow \text{not zero}} u\|^2$$

$$\text{In general, } E(\|y^* - \hat{u}\|^2) = n\sigma^2 + p\sigma^2 + \|(I - P)u\|^2$$

↑
goes up when
adding predictors

↑
goes down when
adding predictors.

Goal: Add predictors that reduce bias
 Don't add unnecessary predictors

How to do this? Find a way to estimate the prediction error

IDENTITY #2: $E[\|y^* - \hat{m}\|^2] = E[RSS] + 2p\sigma^2$

This means $RSS + 2p\sigma^2$ is an unbiased est. of $E[PSS]$

- In particular, this shows that
- 1) RSS is a bad est. of PSS
 - 2) RSS gets worse at est PSS as p gets bigger.

In general, don't know σ^2 , but suppose $\hat{\sigma}^2$ is an unbiased est.

Then $RSS + 2p\hat{\sigma}^2$ is an unbiased est. of $E[PSS]$.

Model Selection with C_p

Consider two models, for y ,

$y \sim X_1$	$ncol(X_1) = p_1$	<u>Model 1</u>
$y \sim X_2$	$ncol(X_2) = p_2$	<u>Model 2</u>

- 1) fit M_1 , get RSS_1
- 2) fit M_2 , get RSS_2
- 3) obtain estimate $\hat{\sigma}^2$ of σ^2 by fitting a model with all available regressors.

Then, prefer M_1 to M_2 if $RSS_1 + 2p_1\hat{\sigma}^2 < RSS_2 + 2p_2\hat{\sigma}^2$

$RSS_j + 2p_j\hat{\sigma}^2$ is the " C_p statistic" for model j .

$$C_p(M_j) = RSS_j + 2p_j\hat{\sigma}^2$$

Sometimes, C_p is defined as $C_p(M_j) = \frac{RSS_j}{\hat{\sigma}^2} + 2p_j - n$

but this doesn't matter for model comparison.