

Model selection [F.3.4]

L12

24.02.2017

SPSE = Expected squared prediction error

$E(Y) = \mu$ is the truth, but we model μ as $\mu = X\beta_M$ and we assume

$$Y = X\beta_M + \epsilon, \quad E(\epsilon) = 0, \quad \text{Cov}(\epsilon) = \sigma^2 I$$

$M \subset \{0, 1, 2, \dots, k\}$ and $|M| = \text{size of model} = \# \text{parameters}$
 $\uparrow$
 $\uparrow$
 $\uparrow$
 $X_1, \dots, X_k$
 Our model is based on a subset of all the available covariates

TRAINING: $i = 1, \dots, n$ our observations, available Y_i, X_i^T
 and $\hat{\beta}_M$ is the estimator from the training set for our model M .

VALIDATION: $j = 1, \dots, T$ new observations, available as Y_j and X_j^T .

$$\sum_{j=1}^T E \left((Y_j - \hat{Y}_{jM})^2 \right) = \text{SPSE}$$

$\uparrow$ new obs $\uparrow$ predicted value based on $\hat{\beta}_M$ and X_j^T

$$\begin{aligned}
 &= \overset{\text{lots of work}}{\dots} = \\
 &n\sigma^2 + \underbrace{|M| \cdot \sigma^2}_{\substack{\text{variance} \\ \text{of} \\ \text{model}}} + \underbrace{\sum_{j=1}^7 \left(\underbrace{E(x_j^T \hat{\beta}_M)}_{\substack{\text{expected} \\ \text{true value} \\ \text{of } \mu_j}} - \mu_j \right)^2}_{\substack{\text{Squared bias} \\ \parallel \\ \text{this can be made} \\ \text{smaller} \\ \text{by including} \\ \text{more variables}}} \\
 &\quad \uparrow \quad \quad \quad \downarrow \quad \quad \quad \downarrow \\
 &\text{variance} \quad \quad \text{variance} \quad \quad \text{Squared bias} \\
 &\text{of} \quad \quad \text{of} \quad \quad \parallel \\
 &\text{errors} \quad \quad \text{model} \quad \quad \text{this can be made} \\
 &\parallel \quad \quad \parallel \quad \quad \text{smaller} \\
 &\text{irreducible} \quad \text{can be made} \quad \text{by including} \\
 &\text{prediction} \quad \text{smaller by} \quad \text{more variables} \\
 &\text{error} \quad \quad \text{choosing} \\
 &\quad \quad \text{fewer variables} \\
 &\quad \quad \nwarrow \quad \quad \nearrow \\
 &\quad \quad \text{bias-variance trade-off}
 \end{aligned}$$

Want to minimize SPSE, but difficult since σ^2 and μ_j is unknown.

PLAN: estimate SPSE and choose the M that minimizes this estimate.

Finding the best model $\leftarrow$ all subsets method

↑
smallest SPSE

1) Have k covariates that might be used

$$Y_i = \beta_0 + \beta_1 X_{1i} + \dots + \beta_k X_{ki} + \epsilon_i$$

How many possible models can I make
(want to have intercept)?

$$2 \cdot 2 \cdot 2 \dots 2 = 2^k \text{ possible models}$$

Fit all possible 2^k models.

2) For $M = \{0, 1, 2, \dots, k\}$ choose the model with the smallest SSE.

Ex: Adrain: $k=7 \Rightarrow$ total $2^7 = 128$ possible models

complexity	searched	best model
$ M = \text{0}$ 1	1	
$ M = \text{1}$ 2:	7	X_4 (A1)
$ M = \text{2}$ 3:	$\binom{7}{2} = 21$	X_1, X_3
$ M = \text{3}$ 4:	$\binom{7}{3} = 35$	X_1, X_2, X_3
$ M = \text{4}$ 5:	$\binom{7}{4} = 35$	X_1, X_2, X_3, X_5
$ M = \text{5}$ 6:	$\binom{7}{5} = 21$	X_1, X_2, X_3, X_5, X_7
$ M = \text{6}$ 7:	$\binom{7}{6} = 7$	$X_1, X_2, X_3, X_4, X_5, X_7$
$ M = \text{7}$ 8:	$\binom{7}{7} = 1$	$X_1, X_2, X_3, X_4, X_5, X_6, X_7$
	<u>24</u>	

3) Now we need to choose between these $k+1$ models found in 2). Which criterion should I use?

$$i) R^2_{adj} = 1 - \frac{\frac{SSE}{(n-p)}}{\frac{SST}{n-1}} \quad |M|=p=k+1$$

Ex: Happiness:

$|M|=1$: only intercept

$|M|=2$: love (3) : 60.5

$|M|=3$: love + work (10) : 66.3

$|M|=4$: love + work + money (13) : **68.4**

$|M|=5$: love + work + sex + money (15) : 67.6

$$ii) \text{ Mallows' } C_p = \frac{SSE}{\hat{\sigma}^2_{Fou}} - n + 2|M|$$

$$vs \quad \hat{SPSSE} = SSE + 2|M| \cdot \hat{\sigma}^2_{Fou} \leftarrow \begin{matrix} \uparrow \\ \text{give some} \\ \text{result.} \end{matrix}$$

$$iii) AIC = n \cdot \ln(\hat{\sigma}^2) + 2(|M|+1)$$

$\uparrow$
 $\frac{SSE}{n-p_{k+1}}$

$$iv) BIC = n \cdot \ln(\hat{\sigma}^2) + \ln(n) (|M|+1)$$

BIC gives more penalty than AIC to large models.

Ex: Happy: BIC best model = love + work

Homework: slide 24 Acid rain

Transformation of response and predictors might improve the fit of the regression model.

The BoxCox transform

$$g_{\lambda}(y) = \begin{cases} \frac{y^{\lambda} - 1}{\lambda} & \lambda \neq 0 \\ \ln(y) & \lambda = 0 \end{cases}$$

↑
class of function

For $Y = X\beta + \epsilon$, $\epsilon \sim N(0, \sigma^2 I)$ the best value of λ is based on maximizing the ^{profile} likelihood

$$l(\lambda) = -\frac{n}{2} \ln\left(\frac{SSE_{\lambda}}{n}\right) - (\lambda - 1) \sum_{i=1}^n \ln y_i$$

SSE_{λ} is the SSE when $g_{\lambda}(Y)$ is the response

R: `boxcox(fit)`, see plot.