

Hypothesis testing about β_j

21.02.2017

L11

Previously: $Y = X\beta + \epsilon$, $\epsilon \sim N_n(0, \sigma^2 I)$

$\hat{\beta}_j$ is the j th element of $\hat{\beta} = (X^T X)^{-1} X^T Y$

$$\hat{\beta} \sim N_p(\beta, (X^T X)^{-1} \sigma^2)$$

$$\hat{\sigma}^2 = \frac{1}{n-p} \hat{\epsilon}^T \hat{\epsilon} = \frac{SSE}{n-p}$$

$$\hat{\epsilon} = Y - \hat{Y} = Y - X\hat{\beta}$$

$$T_j = \frac{\hat{\beta}_j - \beta_j}{\sqrt{C_{jj}} \hat{\sigma}} \sim t_{n-p}$$

$\nwarrow$ $\nearrow$
obs # param. we estimate

j th diagonal element of $(X^T X)^{-1}$

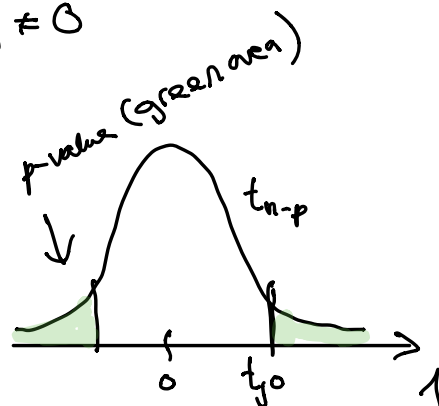
Test for association (linear) between response Y and X_j :

$$H_0: \beta_j = 0 \quad \text{vs} \quad H_a: \beta_j \neq 0$$

When H_0 is true:

$$T_{j0} = \frac{\hat{\beta}_j - 0}{\sqrt{C_{jj}} \hat{\sigma}} \sim t_{n-p}$$

test statistic



$$\begin{aligned}
 \text{p-value: } P(|T_{j0}| \geq |t_{j0}|, H_0 \text{ true}) \\
 = 2 \cdot P(T_{j0} \geq |t_{j0}|)
 \end{aligned}$$

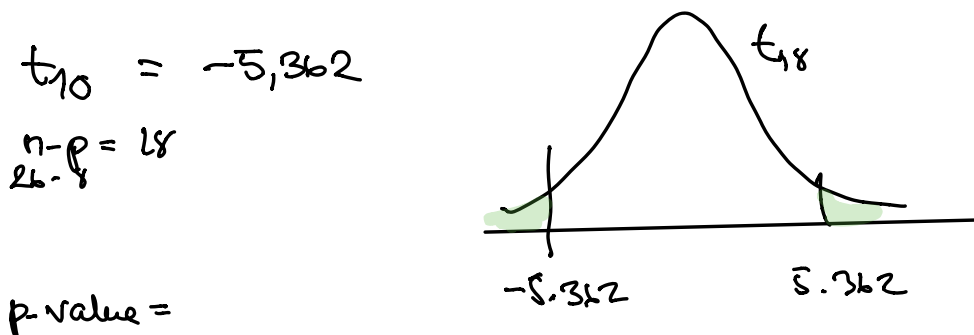
$\uparrow$
 $n-p$

Reject H_0 when $|t_{j0}| > t_{\frac{\alpha}{2}, n-p}$ sign. level α .

Ex: Acid rain: linear association between SO_4 and pH

$$H_0: \beta_1 = 0 \text{ vs } H_1: \beta_1 \neq 0$$

From summary of lmin R (slide)



$$2 \cdot P(T_{18} > 5.362) = \underline{\underline{4.3 \cdot 10^{-5}}}$$

$\uparrow$

$$R: 2 * (1 - pt(5.362, 18))$$

Reject H_0 for all $\alpha > 4.3 \cdot 10^{-5}$.

$\uparrow$
 area to left

Residuals (again)

$$\hat{\epsilon} = Y - \hat{Y}, \quad \hat{\epsilon} \sim N_n(0, \sigma^2(I - H))$$

R: residuals (fit)

The residuals have heteroscedastic variances
 $\text{Var}(\hat{\epsilon}_i) = \sigma^2(1 - h_{ii})$ and $\text{Cov}(\hat{\epsilon}_i, \hat{\epsilon}_j) = \sigma^2(-h_{ij})$

can in general be
 $\neq 0$, but in most
cases experience shows that
 ≈ 0

Standardized residuals:

$$r_i = \frac{\hat{\epsilon}_i}{\hat{\sigma} \sqrt{1 - h_{ii}}}$$

will be (approx.)
homoscedastic.

R: rstandard (fit)

Studentized residuals: fitting the model to
all obs. except i to make r_i^* . ← see slide



use studentized!

R: rstudent (fit)

See example on slide
for r_i vs $\hat{\epsilon}_i$

Analysis of variance decomposition and R^2

$$y_1, \dots, y_n \text{ and } \bar{y} = \frac{1}{n} \sum_{i=1}^n y_i$$

$$\underbrace{\sum_{i=1}^n (y_i - \bar{y})^2}_{\substack{\text{sums of squared} \\ \text{total} \\ \text{SST}}} = \underbrace{\sum_{i=1}^n (y_i - \hat{y}_i + \hat{y}_i - \bar{y})^2}_{\substack{\text{sums of squared} \\ \text{errors} \\ \text{SSE}}} = \dots = \underbrace{\sum_{i=1}^n (\hat{y}_i - \bar{y})^2}_{\substack{\text{sums of sq} \\ \text{regression} \\ \text{SSR}}}$$

↑
explained
by regression

With vectors and matrices:

$$\underbrace{Y^T (I - \frac{1}{n} 11^T) Y}_{\text{SST}} = \underbrace{Y^T (I - H) Y}_{\text{SSE}} + \underbrace{Y^T (H - \frac{1}{n} 11^T) Y}_{\text{SSR}}$$

This is used to define:

$$R^2 = \frac{\text{SSR}}{\text{SST}} = 1 - \frac{\text{SSE}}{\text{SST}}$$

↑
coefficient of
determination

↖ relative proportion of
total variability explained
by the regression

Ex: Acid rain, full model (all covariates available)
 $R^2 = 0.93, 93\%$

Is the regression significant?

$H_0: \beta_1 = \beta_2 = \dots = \beta_k = 0$ vs

H_1 : at least one $\beta_j \neq 0$ $j=1, \dots, k$

Test statistics: $F = \frac{\overset{(SST-SSE)}{\text{SSR}} / k}{SSE / (n-p)} \sim F_{k, n-p}$

↑
 prove this in Part 3 in
 a general setting

Ex: Acid rain:

$$\left. \begin{array}{l} \text{F-observed: } 34.15 \\ k=7, \quad n-p=18 \\ \quad \quad \quad \uparrow \quad \downarrow \\ \quad \quad 25 \quad 8 \end{array} \right\} \underbrace{P(F_{7,18} > 34.15)}_{\text{p-value}} = 3.9 \cdot 10^{-7}$$

Ex: Volume of tree and the lumberjacks

Big model: $Y_i = \beta_0 + \beta_1 X_{1i} + \beta_2 X_{2i} + \beta_3 X_{3i} + \epsilon_i$

Small model: $Y_i = \beta_0 + \beta_1 X_{1i} + \beta_2 X_{2i} + \epsilon_i$

$R^2_{\text{big}} \geq R^2_{\text{small}} \iff$ since $\hat{\beta}_3$ is found

to minimize $\text{SSE} = \hat{\epsilon}^T \hat{\epsilon}$, thus maximize

$$R^2 = 1 - \frac{\text{SSE}}{\text{SST}}$$

$R^2_{\text{big}} = R^2_{\text{small}}$ if $\hat{\beta}_3 = 0$

if $\hat{\beta}_3 \neq 0$ the $\text{SSE}_{\text{big}} < \text{SSE}_{\text{small}}$ and

$R^2_{\text{big}} > R^2_{\text{small}}$.

R^2 will always increase (or stay unchanged) when a new covariate is added to the model.

Next: more on choosing a good model, and then

$$R^2_{\text{adj}} = 1 - \frac{\text{SSE} / (n-p)}{\text{SST} / (n-1)} \leftarrow \begin{array}{l} \text{penalizing} \\ \text{adding} \\ \text{many} \\ \text{covariates} \end{array}$$

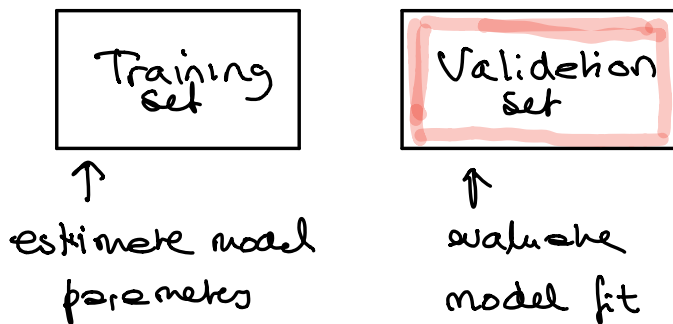
is one criterion to use instead of R^2 for model selection.

Model choice and variable selection [F3.4]

Question 1: Is a full model (all available covariates fitted) the best model?

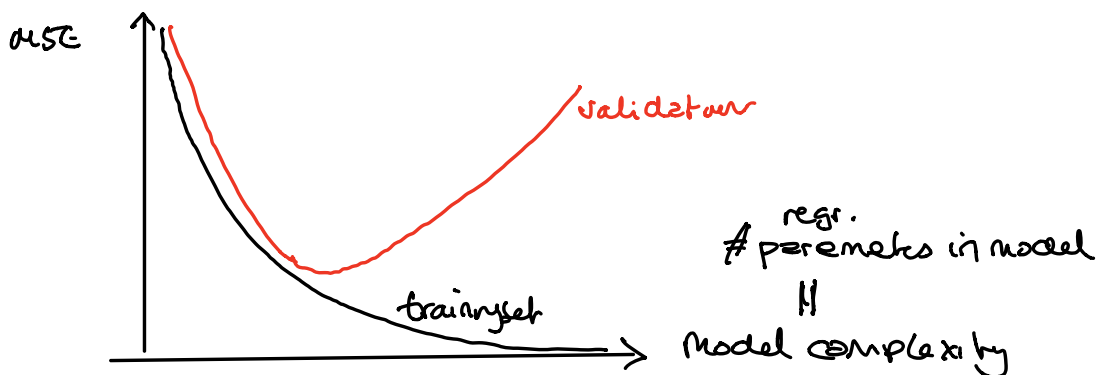
good interpretability good for future predictions

Data is divided into



$$MSE = \frac{1}{n} \sum_{i=1}^n (Y_i - \hat{Y}_i)^2$$

Now calculate the MSE on the training set and on the validation set and plot as a function of model complexity:



Answer 1: No, this may lead to overfitting =
fitting the trend + the noise!

⇒ so, what can we do instead?