

Multivariate

①

$$y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \dots + \beta_{p-1} x_{i,p-1} + \varepsilon_i$$

→ if 5 basic assumptions hold $\Rightarrow$ LS fit OK

$$\begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} = \begin{pmatrix} 1 & x_{11} & x_{12} & \dots & x_{1,p-1} \\ 1 & x_{21} & x_{22} & \dots & x_{2,p-1} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{n1} & x_{n2} & \dots & x_{n,p-1} \end{pmatrix} \begin{pmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_{p-1} \end{pmatrix} + \begin{pmatrix} \varepsilon_1 \\ \vdots \\ \varepsilon_n \end{pmatrix}$$

$$\begin{matrix} \underset{n \times 1}{\underline{y}} = & \underset{n \times p}{\underline{X}} & \underset{p \times 1}{\underline{\beta}} & + & \underset{n \times 1}{\underline{\varepsilon}} \end{matrix}$$

Minimize LS $Q = \sum (y_i - \underline{x}_i' \underline{\beta})^2$ where $\underline{x}_i' = (1, x_{i1}, \dots, x_{i,p-1})$

$$= \sum e_i^2$$

$$= \underline{e}' \underline{e} = \underbrace{(y - \underline{X}\underline{\beta})'(y - \underline{X}\underline{\beta})}_{\text{scalar } |X|}$$

$$\Rightarrow \frac{\partial Q}{\partial \underline{\beta}} = 0$$

$$\Rightarrow \boxed{\text{Normal Equations } (\underline{X}'\underline{X})\underline{\beta} = \underline{X}'\underline{y}}$$

(A) In simple linear regression $\rightarrow$ solution $\hat{\beta}_1 = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_j - \bar{x})^2} = \frac{\text{cov}(x, y)}{\text{var}(x, x)}$

(B) Multiple case $\rightarrow (\underline{X}'\underline{X}) \sim \text{cov}(\underline{X})$ where $v(x_1), v(x_2), \dots, v(x_{p-1})$ on the diagonal, and cross-covariances $\text{cov}(x_1, x_2), \text{cov}(x_u, x_e)$ off-diagonal

$$\rightarrow (\underline{X}'\underline{y}) = \begin{pmatrix} \text{cov}(x_1, y) \\ \text{cov}(x_2, y) \\ \vdots \\ \text{cov}(x_{p-1}, y) \end{pmatrix}$$

(2)

• If $(X'X) = \Lambda$ (diagonal), all x 's are uncorrelated.

$$\rightarrow \text{then } \hat{\beta} = (X'X)^{-1} X'y \sim \begin{pmatrix} \text{corr}(x_1, y) \\ \vdots \\ \text{corr}(x_{p+1}, y) \end{pmatrix}$$

i.e. each coefficient has a direct interpretation as the dependency between an individual x -variable and y .

• If $(X'X)$ not diagonal $\rightarrow$ no direct interpretation

$\rightarrow$ can't separate influence of one x on y from another (correlated) x .

$\rightarrow$ That is, $\hat{\beta}_h$ involves both $\text{corr}(x_h, y)$ and

$\text{corr}(x_l, y), l \neq h$.

$\text{corr}(x_l, x_h)$

• Moreover $\rightarrow$ if x 's are highly correlated, $(X'X)^{-1}$ may not exist! ($\det(X'X) \rightarrow 0$)

or inverse $(X'X)^{-1}$ numerically unstable

$\rightarrow$ When this happens, magnitude & even sign of $\hat{\beta}$'s

may become meaningless. [Many models fit the data equally well]

Ex: $X_2 = a + bX_1$

$$y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \varepsilon$$

} infinite set of β_0, β_1 fit these data equally well

(3)

The Hat-matrix

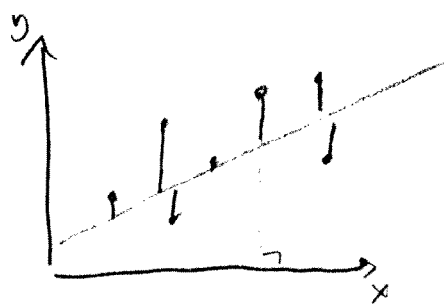
Assume we can solve the normal equations & obtain

$$\hat{\beta} = (X'X)^{-1} X'y$$

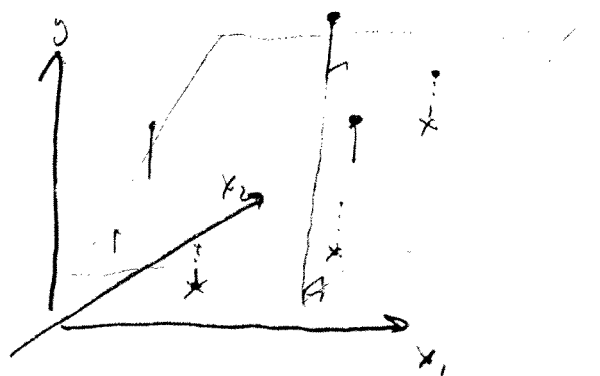
• Fitted values $\hat{y} = X\hat{\beta} = \underbrace{X(X'X)^{-1}X'}_{\text{Hat-matrix}} y$

Hat-matrix $H = X(X'X)^{-1}X'$

- H is a projection matrix — projects the data y onto the plane spanned by X



Simple case



Projection $\perp$ to both x_1 & x_2
in multiple case.

Facts

④

$$\bullet H H = H \quad \begin{cases} (H \text{ is an idempotent matrix}) \\ (H' = H) \\ (\text{Symmetric}) \end{cases} \quad \begin{cases} (\text{No point in redoing regression}) \end{cases}$$

$$\bullet e = \hat{\varepsilon} = y - \hat{y} = y - Hy = (I - H)y$$

$$\bullet X'e = X'(I - X(X'X)^{-1}X')y = 0$$

(orthogonal projection)

$$\bullet \hat{y}'e = y'H(I - H)y = 0$$

(fitted values $\perp$ to residuals)

$$\begin{pmatrix} \hat{y}_1 \\ \hat{y}_2 \\ \vdots \\ \hat{y}_n \end{pmatrix} = \begin{pmatrix} h_{11} & h_{12} & \dots & h_{1n} \\ h_{21} & h_{22} & \dots & h_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ h_{n1} & h_{n2} & \dots & h_{nn} \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix}$$

$\swarrow$ leverage $h_{ii} = X_i (X'X)^{-1} X_i'$

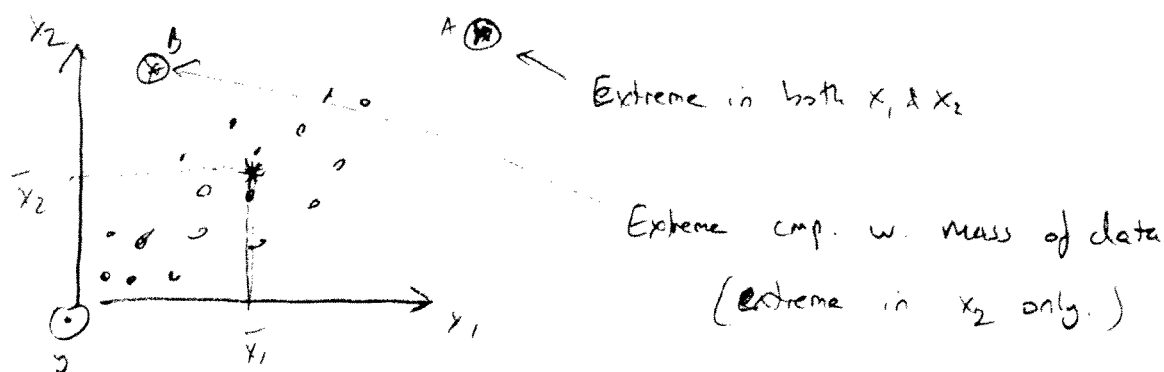
$\swarrow$ i th row of X

Note, h_{ii} is large when X_i is extreme compared with $\bar{X}$

$\swarrow$
vector of all x -variable values for observation i

$\swarrow$
 vector, average x -values.

Extreme how? wrt mass of data



Both $\odot$ observations have high leverage.

Properties

• As before, $E(\hat{\beta}) = \beta$ (unbiased)

$$V(\hat{\beta}) = V(\underbrace{(\mathbf{X}'\mathbf{X})^{-1}}_{\text{Constant}} \underbrace{\mathbf{X}'\mathbf{y}}_{\sim \sigma^2 \mathbf{I}}) = (\mathbf{X}'\mathbf{X})^{-1} \mathbf{X}' V(\mathbf{y}) \mathbf{X} (\mathbf{X}'\mathbf{X})^{-1} = \sigma^2 (\mathbf{X}'\mathbf{X})^{-1}$$

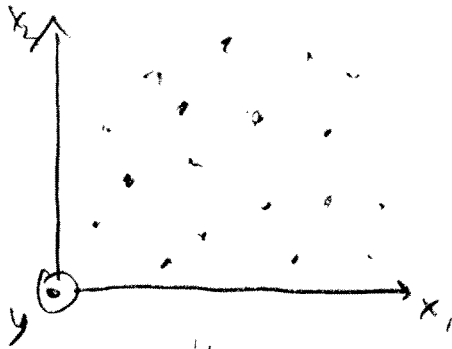
Meaning? $\rightarrow$ More spread in any x ($\text{var}(x_u)$ large)
 $\rightarrow$ better estimate of $\hat{\beta}$.

$\rightarrow \sigma^2 \uparrow \rightarrow \text{variance} \uparrow$

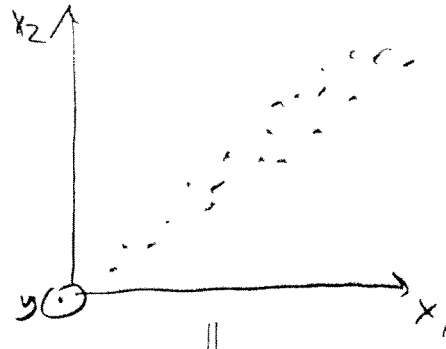
$\rightarrow$ Unless $(\mathbf{X}'\mathbf{X})$ is diagonal, $\hat{\beta}$'s are correlated $\Rightarrow$ Think about impact on inference!

$\rightarrow$ The more dependent x 's are, the larger the $V(\hat{\beta})$.

⑥



↓
 $\hat{\beta}_1$ & $\hat{\beta}_2$ nearly
 uncorrelated



↓
 $\hat{\beta}_1$ & $\hat{\beta}_2$ heavily dependent
 → large variance associated
 with estimates
 → Difficult to make direct
 inferences.

Fitted values $\hat{y} = Hy \Rightarrow E(\hat{y}) = E(y)$

$$V(\hat{y}) = H V(y) H = \sigma^2 H$$

(i.e. large variance @
 locations of high leverage)

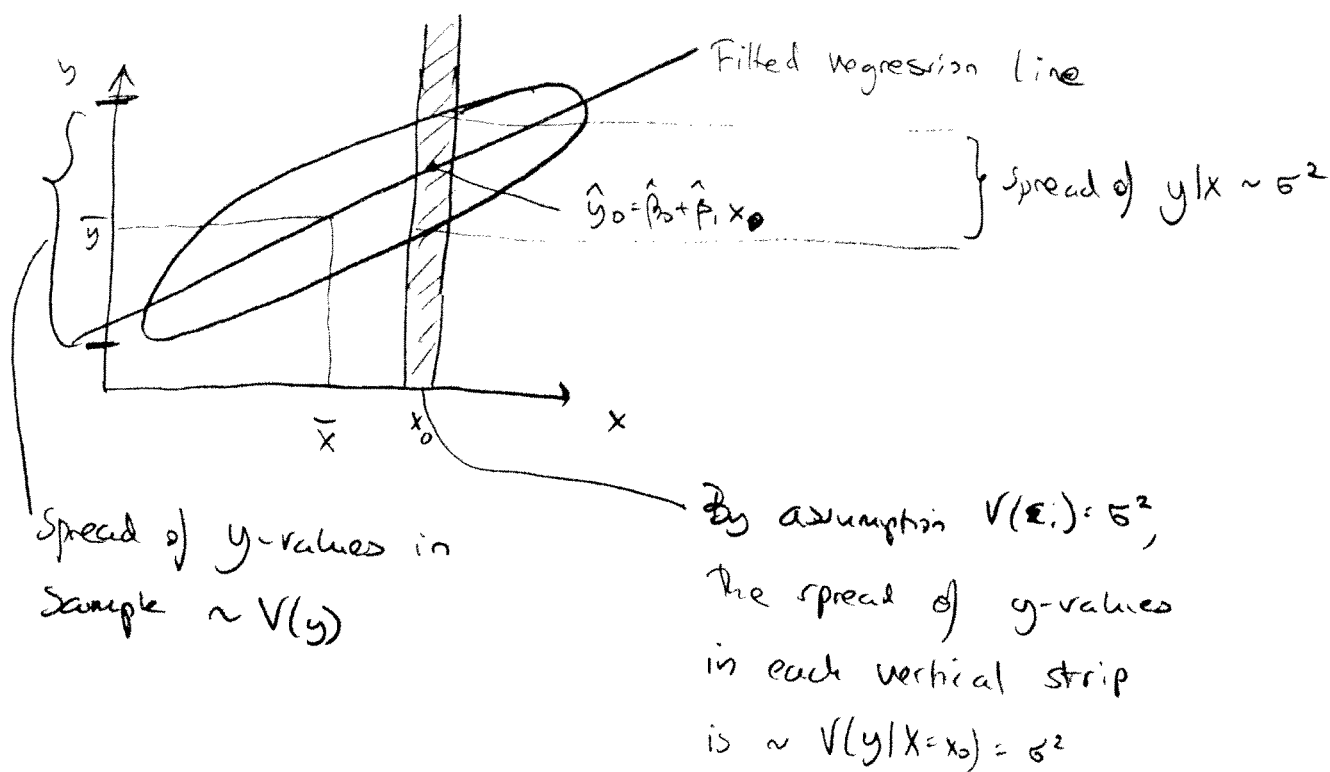
Residuals $e = (I - H)y \Rightarrow E(e) = 0$

$$V(e) = \sigma^2 (I - H)$$

↓ → smaller variance @
 locations w. high leverage
 e 's are correlated!

Variance Decomposition Assume $y_i = \beta_0 + \beta_1 x_i + \varepsilon_i$.

①



- Now, if x and y are unrelated, then $V(y) \approx V(y|x=x)$.
- If $y = \beta_0 + \beta_1 x$ exactly, $V(y|x=x) = 0$!

So how much did linear regression reduce the uncertainty in y ?

① $RSS = SSE = \sum (y_i - \hat{y}_i)^2$ (spread of y 's around the regression line)

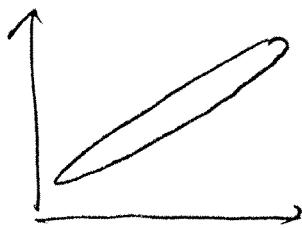
② $SST = \sum (y_i - \bar{y})^2$ (spread around mean of y)

- Since $\hat{y}$ is obtained from the data (the best possible line is fit to the n sample points), RSS is always less than SST

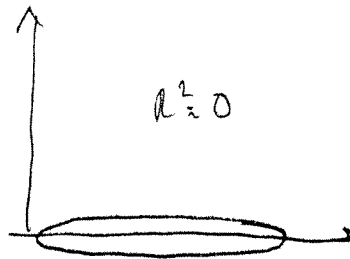
→ Numerical summary; $R^2 = \frac{SST - RSS}{SST} = \frac{\text{reduction in spread}}{\text{original spread}}$

$$= 1 - \frac{RSS}{SST}$$

R^2 : the % of variance (spread) in y explained by the regression line.



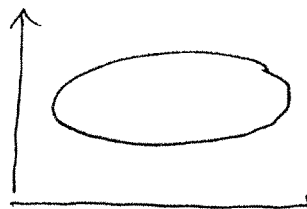
$$R^2 \approx 1$$



$$R^2 \approx 0$$



$$R^2 \approx 0.4$$

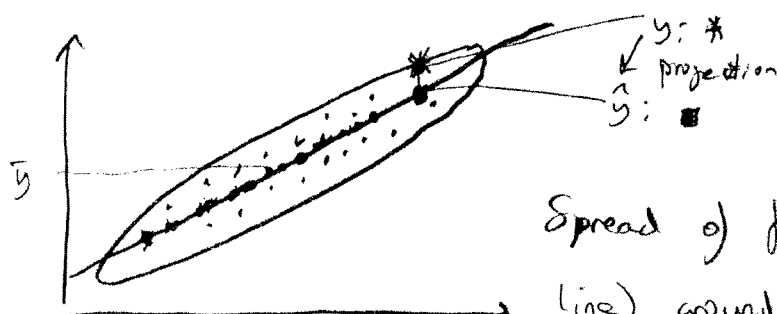


$$R^2 \approx 0$$

Variance decomposition $\Rightarrow$

3

$$\begin{aligned}
 SS_T - SS_E &= \sum (y_i - \bar{y})^2 - \sum (y_i - \hat{y}_i)^2 \\
 &= \sum \cancel{y_i^2} + \sum \bar{y}^2 - 2\bar{y} \sum y_i - \sum \cancel{y_i^2} - \sum \hat{y}_i^2 + 2 \sum y_i \hat{y}_i \\
 &\quad \underbrace{\sum \bar{y}^2 - 2\bar{y} \sum y_i + \sum y_i^2}_{-n\bar{y}^2} \\
 &= -n\bar{y}^2 - \sum \hat{y}_i^2 + 2 \sum y_i \hat{y}_i + \underbrace{2 \sum e_i \hat{y}_i}_{=0 \text{ by } \perp \text{ projection}} \\
 &= \sum \hat{y}_i^2 - n\bar{y}^2 = \sum (\hat{y}_i - \bar{y})^2
 \end{aligned}$$



Spread of fitted values (on the line) around $\bar{y} = \sum (\hat{y}_i - \bar{y})^2 = SS_{reg}$

So, we have

$$SS_T = SS_E + SS_{reg}$$

$$R^2 = 1 - \frac{SS_E}{SS_T} = \frac{SS_{reg}}{SS_T}$$

How large should R^2 be for regression to "make sense"?

prediction
~30%

Interpretation
~80%

'Roughly'

More formally

9

⇒ The F-test Goodness-of-fit test (or lack-of-fit)

Assume y and x are not linearly related, i.e. $\beta_1 = 0$!

Then, what would happen to SS_{reg} & SSE ? (R^2)

- Well, if $\beta_1 = 0$ then true model for y is $y_i = \beta_0 + \varepsilon_i$ ($V(\varepsilon_i) = \sigma^2$)

and our best estimate for β_0 is $\hat{\beta}_0 = \bar{y}$

The expected value of SS_T under the null ($\beta_1 = 0$) is

$$E_0 \left(\sum (y_i - \bar{y})^2 \right) = (n-1)\sigma^2 \quad (\text{just our basic estimate of the variance, remember})$$

- What if $\beta_1 \neq 0$ (or maybe it is, but we fit the line to the data anyway), then

$$\begin{aligned} E \left(\sum (y_i - \hat{y}_i)^2 \right) &= E \left(\sum e_i^2 \right) = \sum_i E(y_i^2) + E(\hat{y}_i^2) - 2E(y_i \hat{y}_i) = \\ &= \left[\begin{array}{l} V(\hat{y}_i) = \sigma^2 h_{ii} \\ \sum e_i \hat{y}_i = 0 \\ \sum h_{ii} = \sum_{i=1}^n \left(\frac{1}{n} + \frac{(x_i - \bar{x})^2}{\sum (x_i - \bar{x})^2} \right) = 2 \end{array} \right] = \sum V(y_i) + E(y_i)^2 + V(\hat{y}_i) + E(\hat{y}_i)^2 \\ &\quad - 2 \underbrace{(E(\hat{y}_i + e_i, \hat{y}_i) - 2E(y_i)E(\hat{y}_i))}_{-2V(\hat{y}_i)} = \end{aligned}$$

$$= \sum V(y_i) - \sum V(\hat{y}_i) = \sigma^2 \cdot n - \sigma^2 \cdot 2 = (n-2)\sigma^2$$

[Note, if we have p parameters in the regression model $E(SSE) = \sigma^2(n-p)$]

• What about SS_{reg} ?

(5)

Under the null / $\beta_1 = 0$, $E(y_i) = \beta_0$

$$\Rightarrow E_0 \left(\sum (y_i - \bar{y})^2 \right) = E_0 \left(\sum \hat{y}_i^2 + \bar{y}^2 - 2\bar{y}\hat{y}_i \right) = E_0(\sum \hat{y}_i^2) + n E_0(\bar{y}^2) - 2 E_0(\bar{y} \sum \hat{y}_i)$$

$$= \underbrace{\sigma^2 \sum h_{ii} + n\beta_0^2}_{(1)} + \underbrace{\sigma^2 + n\beta_0^2}_{(2)} - 2 \underbrace{\sigma^2 - 2n\beta_0^2}_{(3)} =$$

$$\left(\sum V(\hat{y}_i) + \sum E(\hat{y}_i)^2 \right) \left(n V(\bar{y}) + E(\bar{y})^2 \right) - 2 \left(E(\bar{y} \sum \hat{y}_i) = \sum (cov(y_i, \hat{y}_i) + E(y_i)E(\hat{y}_i)) \right)$$

$$= \sigma^2 (\sum h_{ii} - 1) = \sigma^2$$

$$\left[\begin{array}{l} \text{Note, in general w. } p \text{ parameters in} \\ \text{model, } \sum h_{ii} = p \text{ and} \\ E_0(SS_{reg}) = \sigma^2 (p-1) \end{array} \right]$$

MEANING:

(y) $\beta_1 = 0$ (X and y are not related)

all SS 's provide estimates of the same parameter, σ^2 !

Under the null ($\beta_1 = 0$)

$$\frac{SST}{n-1} = \hat{\sigma}^2, \quad \frac{SSE}{n-p} = \hat{\sigma}^2 \quad \text{and} \quad \frac{SS_{reg}}{p-1} = \hat{\sigma}^2$$

Extra assumption: $\epsilon_i \sim N(0, \sigma^2)$

6

Now, i) $\epsilon_i \sim$ normally distributed and $\beta_1 = 0$

$$\frac{SS_{\text{reg}}}{\sigma^2} \sim \chi^2_{p-1} \quad \perp \quad \frac{SSE}{\sigma^2} \sim \chi^2_{n-p}$$

sum of sq's
of normals.

independent ($\perp$ projection)

Fact: $\frac{\chi^2_{m/m}}{\text{indep} - \frac{\chi^2_{n/n}}{n}} \equiv F_{m,n}$

Our test statistic: $F = \frac{SS_{\text{reg}} / (p-1)}{SSE / (n-p)} = \frac{(SS_T - SSE) / ((n-1) - (n-p))}{SSE / (n-p)}$

$= \frac{(\text{reduction in SS})}{(\text{reduction in \# parameters})}$

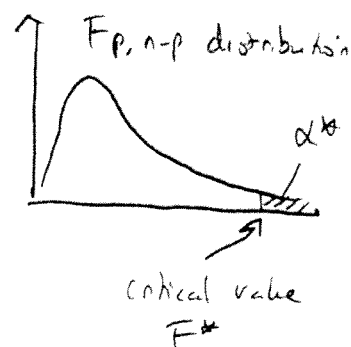
(MSE of full model)

{ (1) $\beta_1 = 0$, $F \sim F_{p-1, n-p}$

(2) $\beta_1 \neq 0$, F is inflated

Hypothesis testing: $\begin{cases} \text{Null: } \beta_1 = 0 \\ \text{Alt: } \beta_1 \neq 0 \end{cases}$

(compare F_{observed} to



We reject $[\beta_1 = 0]$ @ the α^* -level $\Rightarrow$

F_{observed} exceeds F^* .

Discussion

Significant is not the same thing as important!
If n is large enough, we reject almost any hypothesis

Some more distribution theory, assuming $\varepsilon_i \sim N(0, \sigma^2)$

$\hat{\beta}_1 = \sum k_i y_i \rightarrow$ (a) if n is large, by CLT $\hat{\beta}_1 \sim$ normally distributed

$\rightarrow$ (b) if $\varepsilon_i \sim N(0, \sigma^2) \Rightarrow y_i \sim N \Rightarrow \hat{\beta}_1 \sim \dots - \dots - \dots$

(case b) if $\varepsilon_i \sim N(0, \sigma^2) \Rightarrow y_i \sim N(\beta_0 + \beta_1 x_i, \sigma^2)$

$$\Rightarrow \hat{\beta}_1 \sim N\left(\beta_1, \frac{\sigma^2}{\sum (x_i - \bar{x})^2}\right) \quad \text{or} \quad \boxed{\frac{\hat{\beta}_1 - \beta_1}{\sqrt{\frac{\sigma^2}{\sum (x_i - \bar{x})^2}}} \sim N(0, 1)}$$

(but) we don't know $\sigma^2 \Rightarrow$ use plug-in estimator $\hat{\sigma}^2$

From before we know $\frac{SSE}{n-2} = \hat{\sigma}^2$ and if ε_i 's are

normal $\frac{SSE}{\sigma^2} \sim \chi^2_{n-2}$ or equivalently $(n-2) \frac{\hat{\sigma}^2}{\sigma^2} \sim \chi^2_{n-2}$

Fact [review basic stats]

(8)

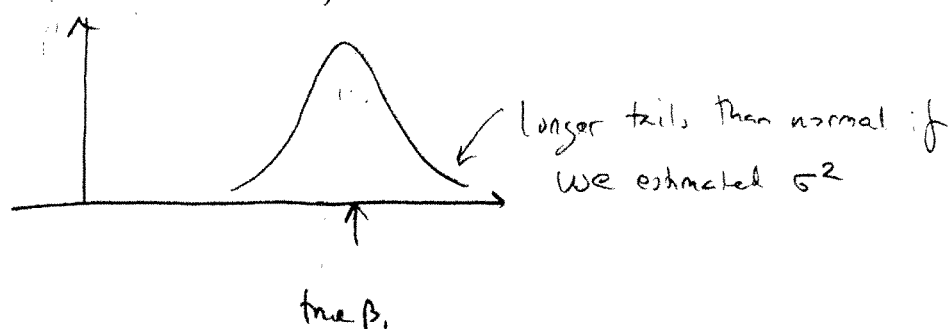
$$\text{index} - \frac{N(0,1)}{\sqrt{\chi^2_{v/r}}} \equiv t_v \quad \Rightarrow \text{here; } \frac{\hat{\beta}_1 - \beta_1}{\sqrt{\hat{\sigma}^2 / \sum (x_j - \bar{x})^2}} \bigg/ \sqrt{\frac{\hat{\sigma}^2}{\sigma^2}}$$

$$= \frac{\hat{\beta}_1 - \beta_1}{\sqrt{\hat{\sigma}^2 / \sum (x_j - \bar{x})^2}} = \frac{\hat{\beta}_1 - \beta_1}{SE(\hat{\beta}_1)} = t\text{-statistic}$$

The standard error

$$\sqrt{\widehat{V(\hat{\beta}_1)}} \leftarrow \sigma^2 \text{ estimated.}$$

So, sampling distribution of estimated $\hat{\beta}_1$.



$\Rightarrow$ Back to Basics now!

• Interval estimate $\left[\hat{\beta}_1 \pm t_{n-p}(1-\alpha/2) \sqrt{\frac{\hat{\sigma}^2}{\sum (x_j - \bar{x})^2}} \right] = I_{\beta_1}$

- I_{β_1} covers true β_1 w. probability $1-\alpha$.

- Question: Does I_{β_1} cover 0 (the null)?