

①

STAT 423

WIN 2016

Properties of OLS Estimators, Part II

Objective: Infer properties of a $\left\{ \begin{array}{l} \text{process} \\ \text{population} \end{array} \right\}$ from data from $\left\{ \begin{array}{l} \text{an experiment} \\ \text{a survey} \end{array} \right\}$

Post-experiment: Data are fixed, known
 $\hat{\beta}_0, \hat{\beta}_1$ are fixed, known

Pre-experiment: Data are unknown, random
 $\hat{\beta}_0, \hat{\beta}_1$ are unknown, random

Pre-experimental Question: Is the probability dist. of $(\hat{\beta}_0, \hat{\beta}_1)$ centered around a meaningful representation of the process or population?

Probability Model of an experiment

Experiment: An experimental outcome y_i is generated for each experimental condition x_i , $i=1, \dots, n$.

Linear Model: $E[Y_i | X_i = x_i] = \beta_0 + \beta_1 x_i$ for some β_0, β_1

equivalently, $\left\{ Y_i = \beta_0 + \beta_1 x_i + \epsilon_i, \quad E[\epsilon_i] = 0 \right\} \quad (A1)$

Return to the question: Are the values of $(\hat{\beta}_0, \hat{\beta}_1)$ we will get from the data likely to be representative of the population or process?

Result 1: If (A1) holds, then $E[\hat{\beta}_0] = \beta_0$
 $E[\hat{\beta}_1] = \beta_1$ $\left((\hat{\beta}_0, \hat{\beta}_1) \text{ are unbiased for } (\beta_0, \beta_1) \right)$

$$\text{Pf: } E[\hat{\beta}_1 | x] = E\left[\frac{S_{xy}}{S_{xx}} | x\right] = \frac{1}{S_{xx}} E\left[\sum (x_i - \bar{x})(y_i - \bar{y}) | x\right]$$

$$= \frac{1}{S_{xx}} \sum (x_i - \bar{x}) E[(y_i - \bar{y}) | x]$$

$$= \frac{1}{S_{xx}} \sum (x_i - \bar{x}) \beta_1 (x_i - \bar{x}) = \beta_1 \frac{S_{xx}}{S_{xx}} = \beta_1$$

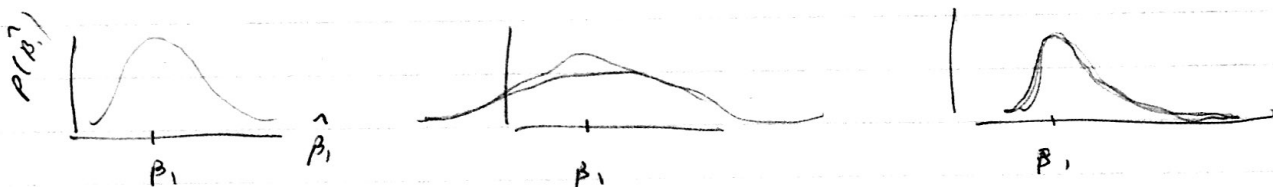
$$\beta_0: \hat{\beta}_0 = \bar{y} - \hat{\beta}_1 \bar{x}$$

$$E[\hat{\beta}_0 | \underline{x}] = E[\bar{y} | \underline{x}] - E[\hat{\beta}_1 | \underline{x}] \bar{x}$$

$$= \beta_0 + \beta_1 \bar{x} - \beta_1 \bar{x} = \beta_0 \quad (\text{for all } \underline{x}).$$

Result 1 does not require $\left\{ \begin{array}{l} \text{independence} \\ \text{normality} \\ \text{constant variance} \end{array} \right.$

but it does not tell us much:



Result #1 doesn't say how concentrated the dist. of $\hat{\beta}_1$ is around β_1 .

Concentration of $(\hat{\beta}_0, \hat{\beta}_1)$ around (β_0, β_1) can be measured with variance and covariance. But to say something about variance, we need to make additional assumptions:

$$(A1) \quad y_i = \beta_0 + \beta_1 x_i + \epsilon_i, \quad E[\epsilon_i | \underline{x}] = 0$$

$$(A2) \quad \begin{array}{l} \text{Var}(\epsilon_i | \underline{x}) = \sigma^2 \text{ for all } i \text{ (constant variance)} \\ \text{Cov}(\epsilon_i, \epsilon_j | \underline{x}) = 0 \text{ for all } i, j \text{ (uncorrelated errors)} \end{array}$$

Result 2: If (A1) + (A2) hold, then

$$E[(\hat{\beta}_0, \hat{\beta}_1)] = (\beta_0, \beta_1) \quad \underline{\text{and}} \quad \text{Var}(\hat{\beta}_1 | \underline{x}) = \sigma^2 \times \frac{1}{S_{xx}}$$

$$\text{Var}(\hat{\beta}_0 | \underline{x}) = \sigma^2 \left(\frac{1}{n} + \bar{x}^2 \frac{1}{S_{xx}} \right)$$

$$\text{Cov}(\hat{\beta}_0, \hat{\beta}_1 | \underline{x}) = -\sigma^2 \left(\frac{\bar{x}}{S_{xx}} \right)$$

$\left\{ \begin{array}{l} \text{will prove} \\ \text{when we} \\ \text{do} \\ \text{multiple} \\ \text{linear} \\ \text{reg.} \end{array} \right.$

3

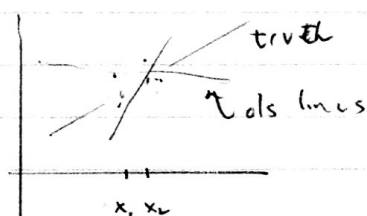
Interpretation

$$\text{Var}(\hat{\beta}_1 | \underline{x}) = \sigma^2 / S_{xx}$$

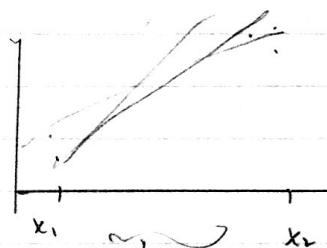
$$= \frac{\sigma^2}{\sum (x_i - \bar{x})^2} = \frac{\sigma^2}{n} \frac{1}{S_x^2}$$

- Depends on $\underline{x}$
- Small if $\left\{ \begin{array}{l} \text{experimental error } \sigma^2 \text{ is small} \\ n \text{ is big} \\ S_x^2 \text{ is big} \end{array} \right. \Rightarrow \text{distribution of } \hat{\beta}_1 \text{ is highly concentrated around } \beta_1$

Picture :



small variance in $\hat{\beta}_1$

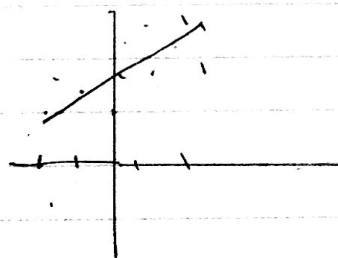


big variance in $\hat{\beta}_1$

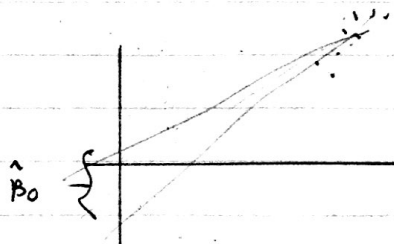
$$\text{Var}(\hat{\beta}_0 | \underline{x}) = \sigma^2 \times \left(\frac{1}{n} + \frac{\bar{x}^2}{S_{xx}} \right)$$

$$= \frac{\sigma^2}{n} \left(1 + \frac{\bar{x}^2}{S_x^2} \right) = \frac{\sigma^2}{n} \frac{1}{S_x^2} (\sum x_i^2)$$

Same interpretation as for $\hat{\beta}_1$, except $\text{Var}(\hat{\beta}_0)$ is larger
 x_i 's are not centered around zero.



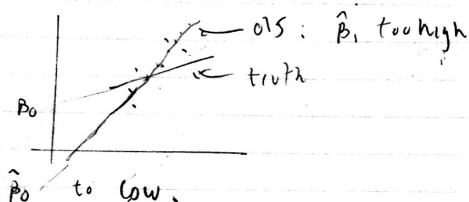
versus



$$\text{Cov}(\hat{\beta}_0, \hat{\beta}_1 | \underline{X}) = -\sigma^2 \left(\frac{\bar{X}}{S_{XX}} \right) \quad \text{Why negative?}$$

Interpretation: if $\hat{\beta}_1$ is higher than β_1 , expect $\hat{\beta}_0$ lower than β_0 .

Picture:



Math:

$$\begin{aligned} \text{Cov}(\hat{\beta}_0, \hat{\beta}_1 | \underline{X}) &= \text{Cov}(\bar{Y} - \hat{\beta}_1 \bar{X}, \hat{\beta}_1 | \underline{X}) \\ &= \text{Cov}(\bar{Y}, \hat{\beta}_1 | \underline{X}) - \text{Cov}(\hat{\beta}_1 \bar{X}, \hat{\beta}_1 | \underline{X}) \\ &= 0 - \bar{X} \text{Cov}(\hat{\beta}_1, \hat{\beta}_1 | \underline{X}) \\ &= -\bar{X} \text{Var}(\hat{\beta}_1 | \underline{X}) \\ &= -\bar{X} \frac{\sigma^2}{S_{XX}} \end{aligned}$$

Exercise: Suppose you want to design an experiment to estimate β_0, β_1 . How should you choose your x_i 's?

Estimating the variance

Q: Is $\hat{\beta}_1$ a precise est. of β_1 ?

A: $V(\hat{\beta}_1 | \underline{X}) = \frac{\sigma^2}{n} \frac{1}{S_{XX}}$

Q: S_{XX} is known, n is known, but σ^2 isn't. So, is $\hat{\beta}_1$ a precise est of β_1 ?

A: Need to know σ^2 .

(5)

Idea: Under (A1) + (A2),

$$y_i = \beta_0 + \beta_1 x_i + \epsilon_i \quad \text{Var}(\epsilon_i | \underline{x}) = \sigma^2, \quad \text{Cov}(\epsilon_i, \epsilon_j | \underline{x}) = 0$$

$$\epsilon_i = y_i - [\beta_0 + \beta_1 x_i]$$

$$\text{Sample var } \epsilon_i \stackrel{?}{\approx} \sigma^2$$

But don't have ϵ_i , as we don't know β_0, β_1 .

Idea: $\hat{\beta}_0 \approx \beta_0, \hat{\beta}_1 \approx \beta_1$, so $\epsilon_i \approx y_i - [\hat{\beta}_0 + \hat{\beta}_1 x_i] = \hat{\epsilon}_i = \underline{\text{residual}}$

$$\text{Var}(\hat{\epsilon}_1, \dots, \hat{\epsilon}_n) \approx \text{Var}(\epsilon_1, \dots, \epsilon_n) \approx \sigma^2$$

More Precisely: let $\text{RSS} = \sum (y_i - [\hat{\beta}_0 + \hat{\beta}_1 x_i])^2 = \sum \hat{\epsilon}_i^2$

$$\text{then } E[\text{RSS}] = \sigma^2 \times (n-2)$$

$$E[\text{RSS}/(n-2)] = \sigma^2$$

so $\hat{\sigma}^2 = \frac{\text{RSS}}{n-2}$ is an unbiased estimate of σ^2 .

Q: Is $\hat{\beta}_1$ a precise est of β_1 ?

$$\underline{A:} \quad \text{V}(\hat{\beta}_1 | \underline{x}) = \frac{\sigma^2}{n} \frac{1}{S_{xx}}$$

$$\text{SD}(\hat{\beta}_1 | \underline{x}) = \frac{\sigma}{\sqrt{n}} \frac{1}{S_{xx}}$$

$$\approx \frac{\hat{\sigma}^2}{n} \frac{1}{S_{xx}}$$

$$\approx \frac{\hat{\sigma}}{\sqrt{n}} \frac{1}{S_{xx}} = \text{SE}(\hat{\beta}_1 | \underline{x})$$

The estimated standard deviation is the standard error.

$$\text{Similarly, } \text{SE}(\hat{\beta}_0 | \underline{x}) = \sqrt{\hat{\sigma}^2 \left(1 + \frac{\bar{x}^2}{S_{xx}}\right)} = \hat{\sigma} \sqrt{1 + \frac{\bar{x}^2}{S_{xx}}}$$

$\hat{\sigma}^2, \text{SE}(\hat{\beta}_0 | \underline{x}), \text{SE}(\hat{\beta}_1 | \underline{x})$ will help us with $\left\{ \begin{array}{l} \text{confidence intervals} \\ \text{hypothesis tests} \\ \text{prediction intervals} \end{array} \right.$

Normal Linear Model

Confidence intervals
Hypothesis tests
Probability statements

} require we know more about the distribution of $(\hat{\beta}_0, \hat{\beta}_1)$ beyond expectations, variances.

Result 3: If $\epsilon_1, \dots, \epsilon_n$ are normally distributed, then
 $\hat{\beta}_0, \hat{\beta}_1$ " " also.

Result 3b: Even if $\epsilon_1, \dots, \epsilon_n$ are not normally distributed,
 $\hat{\beta}_0, \hat{\beta}_1$ are approximately normally distributed for large n .

More specifically,

(A1) $y_i = \beta_0 + \beta_1 x_i + \epsilon_i, E[\epsilon_i] = 0$

(A2) $\text{Var}(\epsilon_i | x) = \sigma^2, \text{Cov}(\epsilon_i, \epsilon_j | x) = 0 \text{ all } i, j$

(A3) ϵ_i is normally distributed

(A23) $\epsilon_1, \dots, \epsilon_n \sim \text{i.i.d. } N(0, \sigma^2)$

Result 3: If A1 + A2 + A3 hold, then ↖ given in result 2.

$$\begin{pmatrix} \hat{\beta}_0 \\ \hat{\beta}_1 \end{pmatrix} \sim N \left(\begin{pmatrix} \beta_0 \\ \beta_1 \end{pmatrix}, \text{Cov} \left(\begin{pmatrix} \hat{\beta}_0 \\ \hat{\beta}_1 \end{pmatrix} \right) \right)$$

This says: $E \left[\begin{pmatrix} \hat{\beta}_0 \\ \hat{\beta}_1 \end{pmatrix} | x \right] = \begin{pmatrix} \beta_0 \\ \beta_1 \end{pmatrix}$ ✓ already showed this

$\text{Cov} \left(\begin{pmatrix} \hat{\beta}_0 \\ \hat{\beta}_1 \end{pmatrix} | x \right) = \text{pcov}$ ✓ " "

$\begin{pmatrix} \hat{\beta}_0 \\ \hat{\beta}_1 \end{pmatrix}$ are jointly normally distributed. * Need to show this.

7

Normality of $\hat{\beta}_1$

Recall: (1) $Z \sim N(0, 1) \Rightarrow a + bZ \sim N(a, b^2)$

(2) $w_1, \dots, w_n$ indep. normal, $w_i \sim N(\mu_i, \sigma_i^2)$

then $\underline{a} \cdot \underline{w} = \sum a_i w_i \sim N(\sum a_i \mu_i, \sum a_i^2 \sigma_i^2)$

$$\hat{\beta}_1 = \frac{S_{xy}}{S_{xx}} = \frac{1}{S_{xx}} \cdot (\underline{x} - \bar{x}\underline{1})^T (\underline{y} - \bar{y}\underline{1})$$

$$(\underline{x} - \bar{x}\underline{1}) = (\underline{I} - \underline{1}\underline{1}^T/n) \underline{x} = \underline{C} \underline{x}$$

$$\Rightarrow (\underline{x} - \bar{x}\underline{1})^T (\underline{y} - \bar{y}\underline{1}) = (\underline{x} - \bar{x}\underline{1})^T \underline{C} \underline{y}$$

$$\Rightarrow \hat{\beta}_1 = \frac{1}{S_{xx}} (\underline{x} - \bar{x}\underline{1})^T \underline{C} \underline{y}$$

$$= \underline{a} \cdot \underline{y}$$

Now by (1), $y_i = (\beta_0 + \beta_1 x_i) + \epsilon_i \sim N(\beta_0 + \beta_1 x_i, \sigma^2)$

(2) $y_1, \dots, y_n$ indep. normal, so $\underline{a} \cdot \underline{y}$ is normal

(already computed the mean and variance)