

Properties of OLS Estimators, Part I

Q: Suppose we are going to (1) gather data $\{(x_i, y_i), i=1 \dots n\}$
 (2) Compute $(\hat{\beta}_0, \hat{\beta}_1)$, OLS estimators

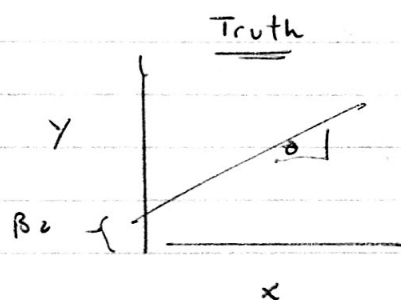
How close will $\{\hat{\beta}_0 + \hat{\beta}_1 x\}$ be to $E[Y|X=x]$?

Two scenarios:

Linear model correct - $E[Y|X=x] = \beta_0 + \beta_1 x$
 for some unknowns (β_0, β_1)

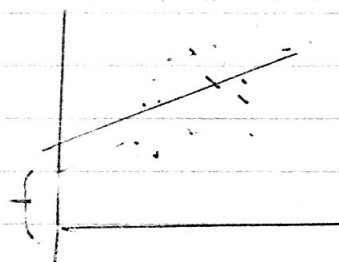
Linear Model incorrect $E[Y|X=x]$ not linear in x .

We first consider the case where the model is correct.

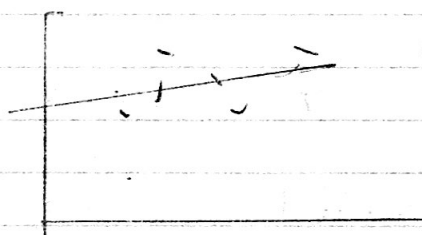


$$\tan^{-1} \theta = \beta_1$$

Dataset 1
OLS Line 1



Dataset 2
OLS Line 2



Etc

The values of $(\hat{\beta}_0, \hat{\beta}_1)$ depend on the outcome of your experiment or study.

If your experiment has random error or noise, or your study is a random sample, then $(\hat{\beta}_0, \hat{\beta}_1)$ are random too!

Truth
 β_0, β_1

Experimental result

$(\hat{\beta}_0, \hat{\beta}_1)_1$
 $(\hat{\beta}_0, \hat{\beta}_1)_2$

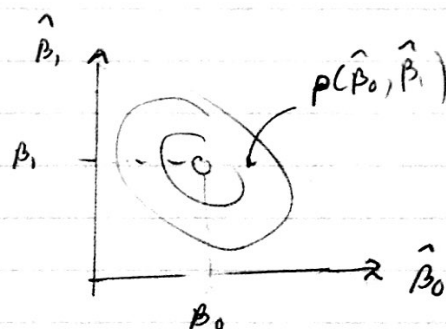
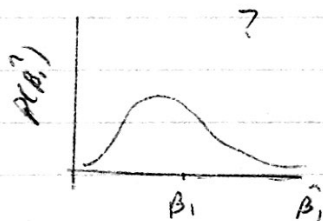
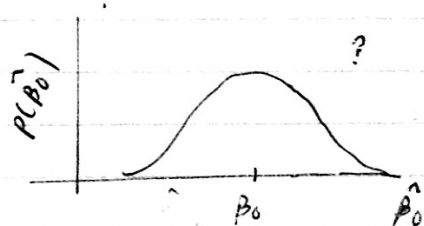
$(\hat{\beta}_0, \hat{\beta}_1)_K$

} different possible experimental results.

The variability of $(\hat{\beta}_0, \hat{\beta}_1)$ across possible experimental outcomes is the sampling variability of $(\hat{\beta}_0, \hat{\beta}_1)$

The probability dist of $(\hat{\beta}_0, \hat{\beta}_1)$ across possible outcomes is the sampling dist. of $(\hat{\beta}_0, \hat{\beta}_1)$.

What do the sampling dist. look like?



Result #1: If $E[Y|X=x] = \beta_0 + \beta_1 x$, then

$$E[\hat{\beta}_0] = \beta_0$$

$$E[\hat{\beta}_1] = \beta_1$$

, i.e. $(\hat{\beta}_0, \hat{\beta}_1)$ are unbiased estimators of (β_0, β_1) .

Proof: Recall
$$\hat{\beta}_1 = \frac{S_{XY}}{S_{XX}} = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2}$$

$$E[\hat{\beta}_1 | \underline{x}] = \frac{\sum (x_i - \bar{x}) E[(y_i - \bar{y}) | \underline{x}]}{\sum (x_i - \bar{x})^2}$$

$$\begin{aligned}
 E[Y_i - \bar{Y} | \underline{x}] &= ? \\
 &= E[Y_i | \underline{x}] - E[\bar{Y} | \underline{x}] \\
 &= (\beta_0 + \beta_1 x_i) - E\left[\frac{1}{n} \sum Y_j \mid \underline{x}\right] \\
 &= (\beta_0 + \beta_1 x_i) - \frac{1}{n} \sum E[Y_j | \underline{x}] \\
 &= (\beta_0 + \beta_1 x_i) - \frac{1}{n} \sum \beta_0 + \beta_1 x_j \\
 &= \beta_0 + \beta_1 x_i - \beta_0 - \beta_1 \bar{x} = \beta_1 (x_i - \bar{x})
 \end{aligned}$$

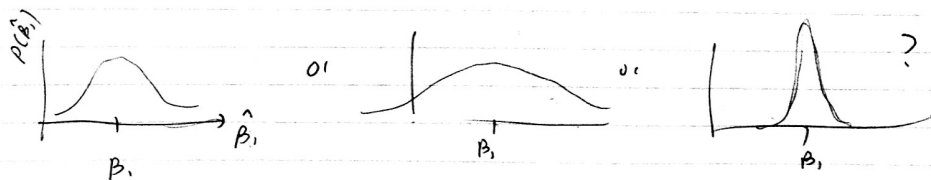
$$\text{So } E[\hat{\beta}_1 | \underline{x}] = \frac{\sum (x_i - \bar{x}) \beta_1 (x_i - \bar{x})}{\sum (x_i - \bar{x})^2} = \beta_1 \left(\frac{\sum (x_i - \bar{x})^2}{\sum (x_i - \bar{x})^2} \right) = \beta_1$$

$$\text{For } \hat{\beta}_0 = \bar{y} - \hat{\beta}_1 \bar{x}$$

$$\begin{aligned}
 E[\hat{\beta}_0 | \underline{x}] &= \beta_0 + \beta_1 \bar{x} - E[\hat{\beta}_1 \bar{x}] \\
 &= \beta_0 + \beta_1 \bar{x} - \beta_1 \bar{x} = \beta_0
 \end{aligned}$$

So: If $E[Y | X=x] = \beta_0 + \beta_1 x$ for some β_0, β_1 ,
 then $E[\hat{\beta}_0, \hat{\beta}_1] = (\beta_0, \beta_1)$.

This says the sampling distributions of $(\hat{\beta}_0, \hat{\beta}_1)$ are centered around the correct values, but it doesn't say how concentrated these distributions are.



Concentration of $(\hat{\beta}_0, \hat{\beta}_1)$ around (β_0, β_1) can be measured with variance and covariance.

- Result #2 if
- ① $E[Y_i | X=x_i] = \beta_0 + \beta_1 x_i$ (linear model)
 - ② $\text{Var}(Y_i | X=x_i) = \sigma^2$ (constant variance)
 - ③ $\text{Cov}(Y_i, Y_{i'} | X_i, X_{i'}) = 0$ (observations are uncorrelated)

then $E[(\hat{\beta}_0, \hat{\beta}_1) | \underline{x}] = (\beta_0, \beta_1)$

and

$$\text{Var}(\hat{\beta}_1 | \underline{x}) = \sigma^2 \times \frac{1}{S_{xx}}$$

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

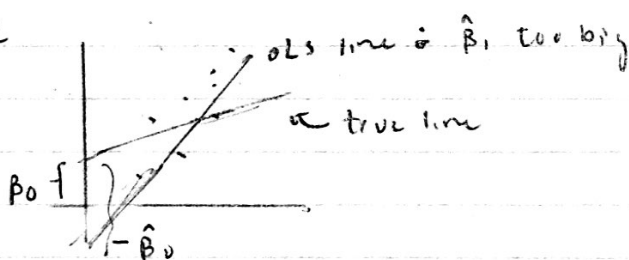
$$\text{Cov}(\hat{\beta}_0, \hat{\beta}_1 | \underline{x}) = -\sigma^2 \left(\frac{\bar{x}}{S_{xx}} \right)$$

} will prove when we get to multiple linear regression.

Why is this negative? What does it mean?

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

Picture



Math:

$$\begin{aligned} \text{Cov}(\hat{\beta}_0, \hat{\beta}_1 | \underline{x}) &= \text{Cov}(\bar{y} - \hat{\beta}_1 \bar{x}, \hat{\beta}_1 | \underline{x}) \\ &= \text{Cov}(\bar{y}, \hat{\beta}_1 | \underline{x}) - \text{Cov}(\hat{\beta}_1 \bar{x}, \hat{\beta}_1 | \underline{x}) \\ &\quad \text{(why 0?)} \\ &= 0 - \bar{x} \text{Cov}(\hat{\beta}_1, \hat{\beta}_1 | \underline{x}) \\ &= -\bar{x} \text{Var}(\hat{\beta}_1 | \underline{x}) \\ &= -\bar{x} \frac{\sigma^2}{S_{xx}} \\ &= \end{aligned}$$