

①

OLS for Multiple Linear Regression

N7

Linear regression model:

$$E[Y_i | X_i] = \beta_0 + \beta_1 X_{i1} + \dots + \beta_p X_{ip}$$

$$= \underline{\beta}^T \underline{x} \quad \left(\underline{\beta} = \begin{pmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_p \end{pmatrix}, \underline{x} = \begin{pmatrix} x_{i1} \\ \vdots \\ x_{ip} \end{pmatrix} \right)$$

$$= \underline{x}^T \underline{\beta}$$

$$= \underline{x} \cdot \underline{\beta}$$

Note: $\underline{x}, \underline{\beta}$ are $(p+1)$ dimensional

Data:

$$\underline{y} = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} \in \mathbb{R}^n$$

$$\underline{X} = \begin{pmatrix} 1 & x_{n1} & \dots & x_{np} \\ \vdots & \vdots & & \vdots \\ 1 & x_{n1} & & x_{np} \end{pmatrix} \in \mathbb{R}^{n \times p}$$

$$= \begin{pmatrix} \underline{x}_1^T \longrightarrow \\ \underline{x}_2^T \longrightarrow \\ \vdots \\ \underline{x}_n^T \longrightarrow \end{pmatrix}$$

Expectation of a vector:

$$\underline{X} \underline{\beta} = \begin{pmatrix} \underline{x}_1^T \longrightarrow \\ \vdots \\ \underline{x}_n^T \longrightarrow \end{pmatrix} \begin{pmatrix} \beta_1 \\ \vdots \\ \beta_p \end{pmatrix} = \begin{pmatrix} \underline{x}_1^T \underline{\beta} \\ \vdots \\ \underline{x}_n^T \underline{\beta} \end{pmatrix}$$

Recall: $E[y_i | x_i] = \underline{x}_i^T \underline{\beta} \Leftrightarrow y_i = \underline{x}_i^T \underline{\beta} + \epsilon_i, E[\epsilon_i | \underline{x}_i] = 0$

Linear Regression model

$$\underline{y} = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} = \begin{pmatrix} \underline{x}_1^T \underline{\beta} + \epsilon_1 \\ \vdots \\ \underline{x}_n^T \underline{\beta} + \epsilon_n \end{pmatrix} = \begin{pmatrix} \underline{x}_1^T \underline{\beta} \\ \vdots \\ \underline{x}_n^T \underline{\beta} \end{pmatrix} + \begin{pmatrix} \epsilon_1 \\ \vdots \\ \epsilon_n \end{pmatrix}$$

$$\underline{y} = \underline{X} \underline{\beta} + \underline{\epsilon}$$

$n \times 1 \quad n \times (p+1) \times 1 \quad n \times 1$

(2)

Linear regression Model (Matrix form)

$$\underline{y} = \underline{X}\underline{\beta} + \underline{\varepsilon} \quad E[\underline{\varepsilon} | \underline{X}] = \underline{0}.$$

Residual Sum of Squares: for observed $\underline{y}$, $\underline{X}$,

$$\begin{aligned} \text{RSS}(\underline{\beta}) &= \sum (y_i - x_i^T \underline{\beta})^2 \\ &= (\underline{y} - \underline{X}\underline{\beta})^T (\underline{y} - \underline{X}\underline{\beta}) \\ &= (\underline{y} - \underline{X}\underline{\beta})^T (\underline{y} - \underline{X}\underline{\beta}) = \|\underline{y} - \underline{X}\underline{\beta}\|^2 \end{aligned}$$

OLS Estimator

The OLS estimator $\hat{\underline{\beta}}$ of $\underline{\beta}$ are the values of $\underline{\beta}$ minimizing $\text{RSS}(\underline{\beta})$

$$\hat{\underline{\beta}} = \arg \min_{\underline{\beta}} \text{RSS}(\underline{\beta})$$

Derivation:

$$\begin{aligned} \text{RSS}(\underline{\beta}) &= (\underline{y} - \underline{X}\underline{\beta})^T (\underline{y} - \underline{X}\underline{\beta}) \\ &= \underline{y}^T \underline{y} - \underline{\beta}^T \underline{X}^T \underline{y} - \underline{y}^T \underline{X} \underline{\beta} + \underline{\beta}^T \underline{X}^T \underline{X} \underline{\beta} \\ &= \underline{y}^T \underline{y} - 2 \underline{\beta}^T \underline{X}^T \underline{y} + \underline{\beta}^T \underline{X}^T \underline{X} \underline{\beta} \end{aligned}$$

$$\frac{\partial \text{RSS}(\underline{\beta})}{\partial \underline{\beta}} = -2 \underline{X}^T \underline{y} + 2 \underline{X}^T \underline{X} \underline{\beta}$$

$$\frac{\partial \text{RSS}(\underline{\beta})}{\partial \underline{\beta}} \Big|_{\hat{\underline{\beta}}} = -2 \underline{X}^T \underline{y} + 2 \underline{X}^T \underline{X} \hat{\underline{\beta}} = 0$$

$$\underline{X}^T \underline{X} \hat{\underline{\beta}} = \underline{X}^T \underline{y}$$

$$\hat{\underline{\beta}} = \frac{\underline{X}^T \underline{y}}{\underline{X}^T \underline{X}} = \underline{(X^T X)^{-1} X^T y}$$

Exercise: Confirm that if $p=1$, result from SLR holds.

(3)

Properties of OLS estimators

(A1): $y = X\beta + \underline{e}$, $E[\underline{e} | X] = 0$, for some $\beta \in \mathbb{R}^{p+1}$

(R1):
$$\begin{aligned} E[\hat{\beta} | X] &= E[(X^T X)^{-1} X^T y | X] \\ &= (X^T X)^{-1} X^T E[y | X] \\ &= (X^T X)^{-1} X^T X \beta = \beta \Rightarrow \hat{\beta} \text{ is unbiased est. of } \beta. \end{aligned}$$

(A2): $e_1, \dots, e_n \sim \text{i.i.d. } N(0, \sigma^2)$ or equivalently,

$\underline{e} \sim N(\underline{0}, \sigma^2 I)$

(R2)
$$\text{Var}(\hat{\beta} | X) \equiv \begin{bmatrix} v(\hat{\beta}_0) & c(\hat{\beta}_0, \hat{\beta}_1) & & \\ c(\hat{\beta}_0, \hat{\beta}_1) & v(\hat{\beta}_1) & & \\ \vdots & & \ddots & \\ \vdots & & & v(\hat{\beta}_p) \end{bmatrix} = \sigma^2 (X^T X)^{-1}$$

(compare to $v(\hat{\beta}_1) = \sigma^2 / s_{xx} = \sigma^2 (s_{xx})^{-1}$ for simple linear regression)

Proof of (R2)

(Lemma 1) (a) If $E[\underline{z}] = 0$, then $\text{Cov}(\underline{z}) = E[\underline{z} \underline{z}^T]$

In general, $\text{Cov}(\underline{w}) = E[(\underline{w} - \underline{\mu})(\underline{w} - \underline{\mu})^T]$, $E(\underline{w}) = \underline{\mu}$

(b) Therefore, if $\underline{w} = \underline{\mu} + \underline{z}$, $E(\underline{z}) = 0$
 $\Rightarrow \text{Cov}(\underline{w}) = \text{Cov}(\underline{z}) = E(\underline{z} \underline{z}^T)$

$\Rightarrow \text{Cov}(\underline{y}) = \text{Cov}(\underline{e}) = E(\underline{e} \underline{e}^T) = \sigma^2 I$

Lemma 2: Let $\underline{y} = \underline{A}(\underline{u} + \underline{z})$, where $E(\underline{z}) = 0$

then $\underline{y} = \underline{A}\underline{u} + \underline{A}\underline{z}$

$$= \underline{u}^* + \underline{z}^* \quad \text{where } \underline{z}^* = \underline{A}\underline{z}$$

$$E(\underline{z}^*) = E(\underline{A}\underline{z}) = \underline{A}E(\underline{z}) = 0$$

then $\text{Var}(\underline{y}) = E[\underline{z}^* \underline{z}^{*T}]$

$$= \underline{A} E(\underline{z}\underline{z}^T) \underline{A}^T = \underline{A} \text{Cov}(\underline{z}) \underline{A}^T$$

Back to regression: $\hat{\beta} = (\underline{X}^T \underline{X})^{-1} \underline{X}^T \underline{y}$

$$= (\underline{X}^T \underline{X})^{-1} \underline{X}^T (\underline{X}\beta + \underline{\epsilon})$$

$$= \beta + (\underline{X}^T \underline{X})^{-1} \underline{X}^T \underline{\epsilon}$$

$$\text{Cov}(\hat{\beta}) = \text{Cov}((\underline{X}^T \underline{X})^{-1} \underline{X}^T \underline{\epsilon})$$

$$= E[(\underline{X}^T \underline{X})^{-1} \underline{X}^T \underline{\epsilon} \underline{\epsilon}^T \underline{X} (\underline{X}^T \underline{X})^{-1}]$$

$$= (\underline{X}^T \underline{X})^{-1} \underline{X}^T E[\underline{\epsilon} \underline{\epsilon}^T] \underline{X} (\underline{X}^T \underline{X})^{-1}$$

$$= (\underline{X}^T \underline{X})^{-1} \underline{X}^T \underline{X} (\underline{X}^T \underline{X})^{-1} \sigma^2 = \underline{\underline{\sigma^2 (\underline{X}^T \underline{X})^{-1}}}$$