

ECON3389 Econometric Methods

Module 1 Multivariate Linear Regression

Alberto Cappello

Department of Economics, Boston College

Spring 2025

Estimation in MLR

- We now have p explanatory variables $X_1, X_2, \dots, X_p$:

$$Y = \beta_0 \cdot 1 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_p X_p + \epsilon$$

- Our sample is $\{y_i, x_{1i}, x_{2i}, \dots, x_{pi}\}_{i=1}^n$. Given estimates $\hat{\beta}_0, \hat{\beta}_1, \dots, \hat{\beta}_p$, our predicted (estimated) outcome is

$$\hat{y}_i = \hat{\beta}_0 + \hat{\beta}_1 x_{1i} + \hat{\beta}_2 x_{2i} + \dots + \hat{\beta}_p x_{pi}$$

- The same logic as in SLR leads us to OLS estimates as the ones that minimize RSS:

$$RSS(\hat{\beta}_0, \hat{\beta}_1, \dots, \hat{\beta}_p) = \sum_{i=1}^n (y_i - \hat{y}_i)^2 \rightarrow \min_{\hat{\beta}_0, \hat{\beta}_1, \dots, \hat{\beta}_p}$$

- Closed form solutions still exist, but they become cumbersome when usual scalar notation is used.

Matrix Notation

$$\mathbf{Y} = \begin{pmatrix} y_1 \\ y_2 \\ \dots \\ y_n \end{pmatrix} \quad \mathbf{X} = \begin{pmatrix} 1 & x_{11} & x_{21} & \dots & x_{p1} \\ 1 & x_{12} & x_{22} & \dots & x_{p2} \\ \dots & \dots & \dots & \dots & \dots \\ 1 & x_{1n} & x_{2n} & \dots & x_{pn} \end{pmatrix} \quad \boldsymbol{\beta} = \begin{pmatrix} \beta_0 \\ \beta_1 \\ \dots \\ \beta_p \end{pmatrix} \quad \boldsymbol{\epsilon} = \begin{pmatrix} \epsilon_0 \\ \epsilon_1 \\ \dots \\ \epsilon_n \end{pmatrix}$$

Matrix Notation

$$\mathbf{Y} = \begin{pmatrix} y_1 \\ y_2 \\ \dots \\ y_n \end{pmatrix} \quad \mathbf{X} = \begin{pmatrix} 1 & x_{11} & x_{21} & \dots & x_{p1} \\ 1 & x_{12} & x_{22} & \dots & x_{p2} \\ \dots & \dots & \dots & \dots & \dots \\ 1 & x_{1n} & x_{2n} & \dots & x_{pn} \end{pmatrix} \quad \boldsymbol{\beta} = \begin{pmatrix} \beta_0 \\ \beta_1 \\ \dots \\ \beta_p \end{pmatrix} \quad \boldsymbol{\epsilon} = \begin{pmatrix} \epsilon_0 \\ \epsilon_1 \\ \dots \\ \epsilon_n \end{pmatrix}$$

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon} \quad \hat{\mathbf{Y}} = \mathbf{X}\hat{\boldsymbol{\beta}} \quad \mathbf{e} = \mathbf{Y} - \hat{\mathbf{Y}}$$

Matrix Notation

$$\mathbf{Y} = \begin{pmatrix} y_1 \\ y_2 \\ \dots \\ y_n \end{pmatrix} \quad \mathbf{X} = \begin{pmatrix} 1 & x_{11} & x_{21} & \dots & x_{p1} \\ 1 & x_{12} & x_{22} & \dots & x_{p2} \\ \dots & \dots & \dots & \dots & \dots \\ 1 & x_{1n} & x_{2n} & \dots & x_{pn} \end{pmatrix} \quad \boldsymbol{\beta} = \begin{pmatrix} \beta_0 \\ \beta_1 \\ \dots \\ \beta_p \end{pmatrix} \quad \boldsymbol{\epsilon} = \begin{pmatrix} \epsilon_0 \\ \epsilon_1 \\ \dots \\ \epsilon_n \end{pmatrix}$$

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon} \quad \hat{\mathbf{Y}} = \mathbf{X}\hat{\boldsymbol{\beta}} \quad \mathbf{e} = \mathbf{Y} - \hat{\mathbf{Y}}$$

Then $RSS = \mathbf{e}'\mathbf{e}$ and our OLS estimates are defined as

$$\hat{\boldsymbol{\beta}}_{OLS} = \underset{\hat{\boldsymbol{\beta}}}{\operatorname{argmin}} RSS = (\mathbf{X}'\mathbf{X})^{-1} \mathbf{X}'\mathbf{Y}$$

Small sample properties

- ZCM assumption ensures that OLS is unbiased:

$$\begin{aligned}
 \mathbb{E}[\hat{\beta}_{OLS}|\mathbf{X}] &= \mathbb{E}\left[(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{Y}|\mathbf{X}\right] = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbb{E}[\mathbf{Y}|\mathbf{X}] = \\
 &= (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbb{E}[\mathbf{X}\beta + \epsilon|\mathbf{X}] = \underbrace{(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{X}\beta}_{=\mathbb{I}} + (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\underbrace{\mathbb{E}[\epsilon|\mathbf{X}]}_{=0} = \\
 &= \beta
 \end{aligned}$$

Small sample properties

- ZCM assumption ensures that OLS is unbiased:

$$\begin{aligned}\mathbb{E}[\hat{\beta}_{OLS}|\mathbf{X}] &= \mathbb{E}\left[(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{Y}|\mathbf{X}\right] = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbb{E}[\mathbf{Y}|\mathbf{X}] = \\ &= (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbb{E}[\mathbf{X}\beta + \epsilon|\mathbf{X}] = \underbrace{(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{X}\beta}_{=\mathbb{I}} + (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\underbrace{\mathbb{E}[\epsilon|\mathbf{X}]}_{=0} = \\ &= \beta\end{aligned}$$

- If ϵ is homoscedastic with no serial correlation, then the variance of OLS is

$$\begin{aligned}\text{Var}[\hat{\beta}_{OLS}|\mathbf{X}] &= \text{Var}\left[\beta + (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\epsilon|\mathbf{X}\right] = \text{Var}\left[(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\epsilon|\mathbf{X}\right] = \\ &= (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\text{Var}[\epsilon|\mathbf{X}]\left((\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\right)' = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'(\sigma^2\mathbb{I})\mathbf{X}(\mathbf{X}'\mathbf{X})^{-1} = \\ &= \sigma^2(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{X}(\mathbf{X}'\mathbf{X})^{-1} = \sigma^2(\mathbf{X}'\mathbf{X})^{-1}\end{aligned}$$

Small Sample Properties

- Just like in SLR, under assumptions of $\mathbb{E}[\epsilon|\mathbf{X}] = 0$ and $\text{Var}[\epsilon|\mathbf{X}] = \sigma^2\mathbb{I}$ OLS is BLUE — has least variance among all linear unbiased estimators.

Small Sample Properties

- Just like in SLR, under assumptions of $\mathbb{E}[\epsilon|\mathbf{X}] = 0$ and $\text{Var}[\epsilon|\mathbf{X}] = \sigma^2\mathbb{I}$ OLS is BLUE — has least variance among all linear unbiased estimators.
- If one assumes normality of the error term ϵ , OLS estimates have exact normal sampling distributions:

$$\hat{\beta}_{OLS} \sim \mathcal{N}\left(\beta, \sigma^2 (\mathbf{X}'\mathbf{X})^{-1}\right)$$

$$\frac{\hat{\beta}_{OLS} - \beta_{OLS}}{\text{se}(\hat{\beta}_{OLS})} \sim t_{n-p-1}$$

- However, assuming normality of the error term is often just as unrealistic in MLR as in SLR. In addition, unbiasedness only works with repeated samples, while we usually have access only to a single sample.

Large Sample Properties (Asymptotics)

Good news — just like in SLR, OLS estimates in MLR are consistent

$$\begin{aligned}\text{plim } \hat{\beta}_{OLS} &= \text{plim } \left((\mathbf{X}'\mathbf{X})^{-1} \mathbf{X}'\mathbf{Y} \right) = \text{plim } \left(\beta + (\mathbf{X}'\mathbf{X})^{-1} \mathbf{X}'\epsilon \right) = \\ &= \beta + \text{plim } \left(\left(\frac{1}{n} \mathbf{X}'\mathbf{X} \right)^{-1} \frac{1}{n} \mathbf{X}'\epsilon \right) = \beta + \text{plim } \left(\left(\frac{1}{n} \mathbf{X}'\mathbf{X} \right)^{-1} \right) \cdot \text{plim } \left(\frac{1}{n} \mathbf{X}'\epsilon \right) = \\ &= \beta + \mathbb{E}[X'X] \cdot \mathbb{E}[X\epsilon] = \beta\end{aligned}$$

and asymptotically normal

Large Sample Properties (Asymptotics)

Good news — just like in SLR, OLS estimates in MLR are consistent

$$\begin{aligned}\text{plim } \hat{\beta}_{OLS} &= \text{plim } \left((\mathbf{X}'\mathbf{X})^{-1} \mathbf{X}'\mathbf{Y} \right) = \text{plim } \left(\beta + (\mathbf{X}'\mathbf{X})^{-1} \mathbf{X}'\epsilon \right) = \\ &= \beta + \text{plim } \left(\left(\frac{1}{n} \mathbf{X}'\mathbf{X} \right)^{-1} \frac{1}{n} \mathbf{X}'\epsilon \right) = \beta + \text{plim } \left(\left(\frac{1}{n} \mathbf{X}'\mathbf{X} \right)^{-1} \right) \cdot \text{plim } \left(\frac{1}{n} \mathbf{X}'\epsilon \right) = \\ &= \beta + \mathbb{E}[\mathbf{X}'\mathbf{X}] \cdot \mathbb{E}[\mathbf{X}\epsilon] = \beta\end{aligned}$$

and asymptotically normal

$$\begin{aligned}\hat{\beta}_{OLS} &\underset{n \rightarrow \infty}{\overset{d}{\rightsquigarrow}} \mathcal{N} \left(\beta, \frac{\sigma^2}{n} (\mathbf{X}'\mathbf{X})^{-1} \right) \\ \sqrt{n} \left(\hat{\beta}_{OLS} - \beta_{OLS} \right) &\underset{n \rightarrow \infty}{\overset{d}{\rightsquigarrow}} \mathcal{N} \left(0, \sigma^2 (\mathbf{X}'\mathbf{X})^{-1} \right)\end{aligned}$$

Statistical Inference

- Standard t-test for significance can interpreted in the same way as in SLR
- Results from example with advertising data, using all three variables:

Variable	Coefficient	SE	t	p-value
Intercept	2.939	0.3119	9.42	<0.0001
TV	0.046	0.0014	32.81	<0.0001
radio	0.189	0.0086	21.89	<0.0001
newspaper	-0.001	0.0059	-0.18	<0.8599

- Correlation matrix:

Variable	TV	radio	newspaper	sales
TV	1.0000	0.0548	0.0567	0.7822
radio		1.0000	0.3541	0.5762
newspaper			1.0000	0.2283
sales				1.00000

Statistical Inference

MLR offers much wider range of possible analysis avenues:

- Is at least one of the predictors $X_1, X_2, \dots, X_p$ useful in predicting the response?
- Do all the predictors help to explain Y , or is only a subset of the predictors useful?
- How well does the model fit the data?
- Given a set of predictor values, what response value should we predict, and how accurate is our prediction?

F-test for linear restrictions

- Standard significance tests only look at one variable at a time. What if we want to assess the *joint significance* of several variables at once?
- This means imposing multiple restrictions at once, e.g. $k = 3$ restrictions:

$$H_0 : \beta_1 = \beta_4 = \beta_5 = 0$$

F-test for linear restrictions

- Standard significance tests only look at one variable at a time. What if we want to assess the *joint significance* of several variables at once?
- This means imposing multiple restrictions at once, e.g. $k = 3$ restrictions:

$$H_0 : \beta_1 = \beta_4 = \beta_5 = 0$$

- The idea of an *F-test* is to compare how much worse our model's fit gets once we impose those restrictions and run OLS with them:

$$F = \frac{(RSS_r - RSS_{ur})/k}{RSS_{ur}/(n - k - 1)} = \frac{(R_{ur}^2 - R_r^2)/k}{(1 - R_{ur}^2)/(n - p - 1)} \sim \mathcal{F}_{k, n-p-1}$$

F-test for linear restrictions

- Standard significance tests only look at one variable at a time. What if we want to assess the *joint significance* of several variables at once?
- This means imposing multiple restrictions at once, e.g. $k = 3$ restrictions:

$$H_0 : \beta_1 = \beta_4 = \beta_5 = 0$$

- The idea of an *F-test* is to compare how much worse our model's fit gets once we impose those restrictions and run OLS with them:

$$F = \frac{(RSS_r - RSS_{ur})/k}{RSS_{ur}/(n - k - 1)} = \frac{(R_{ur}^2 - R_r^2)/k}{(1 - R_{ur}^2)/(n - p - 1)} \sim \mathcal{F}_{k, n-p-1}$$

- If $F > \mathcal{F}_{k, n-p-1}^\alpha$, we reject H_0 on significance level α .

Is at least one predictor useful?

- This question corresponds to special case of

$$H_0 : \beta_1 = \beta_2 = \dots = \beta_p = 0$$

- In Econometrics this is known as regression significance test, and it is automatically performed for every linear regression.
- For our advertising data the result of this test is

$$R^2 = 0.897, \quad F = 570, \quad \text{p-value} < 0.00001 \Rightarrow H_0 \text{ is rejected}$$

Which variables are important?

- Generally we don't have any knowledge as to which variables we should test for joint significance

Which variables are important?

- Generally we don't have any knowledge as to which variables we should test for joint significance
- We have lots of independent variables. We need to decide which is the best model? For our current purposes best = highest Adj R squared

Which variables are important?

- Generally we don't have any knowledge as to which variables we should test for joint significance
- We have lots of independent variables. We need to decide which is the best model? For our current purposes best = highest Adj R squared
- We have to decide two things: The number of variables to include in our model and which variables are those?

Which variables are important?

- Generally we don't have any knowledge as to which variables we should test for joint significance
- We have lots of independent variables. We need to decide which is the best model? For our current purposes best = highest Adj R squared
- We have to decide two things: The number of variables to include in our model and which variables are those?
- The most direct approach is called *best subsets* regression: we compute OLS fit for all possible subsets and then choose between them based on some criterion that balances training error with model size.
 - But this is often not feasible, since they are 2^p possible subsets of p regressors, e.g. with $p = 40$ there are over a billion models!
- Instead, the two most commonly used approaches are *forward selection* and *backward selection*.
- We can also use statistics like Mallows's C_p , Akaike information criterion (AIC), Bayesian information criterion (BIC), Cross-validation (CV)