

First Year Econometrics

Charles Christopher Hyland

July 13, 2026

Contents

1	Probability Theory	6
1.1	Axioms	6
1.2	Distribution Functions and Densities	7
2	Transformations and Expectations	11
2.1	Transformations of Distributions	11
2.2	Expectations of Distributions	12
2.3	Multivariate Normal Distributions	14
2.4	Normal Sampling Distributions	15
2.4.1	Chi-Squared Distribution	15
2.4.2	T and F Distributions	18
3	Linear Regression Model and Linear Projection Model	20
3.1	Conditional Expectation Function	20
3.2	Linear Regression Model	23
3.3	Linear Projection Model	23
4	Regression and Matrix Algebra	26
4.1	Ordinary Least Squares	28
4.2	Partition Regression	33
4.3	Frisch-Waugh-Lovell Theorem	34
4.4	Residual Variance Decomposition	35
5	Asymptotic Theory	37
5.1	Convergence in Distribution and Probability	37
5.1.1	Weak Law of Large Numbers	38
5.1.2	Continuous Mapping Theorem	40
5.1.3	Central Limit Theorem	42
5.1.4	Delta Method	44
5.2	Stochastic Order	46

6	Maximum Likelihood Estimation	51
6.1	Introduction and Set Up	51
6.2	Score Function and Fisher Information Matrix	52
6.3	MLE for Regression	57
6.3.1	Homoskedastic Gaussian Regression	57
6.4	Asymptotic Theory	59
6.5	Hypothesis Testing	60
6.5.1	Likelihood Ratio Test	60
6.5.2	Wald Test	61
6.5.3	Lagrange Multiplier Test	61
7	The General Regression Model	62
7.1	Sample and Population Models	62
7.2	General Regression Model	64
7.3	Classical Linear Regression Models	66
7.4	Finite Properties in Regression Models	68
7.5	Interpretation of Coefficients in Regression Models	70
7.6	Inference in Regression Models	71
7.6.1	Classical Linear Regression Model with Normally Distributed Error Terms	71
7.6.2	Hypothesis Testing	78
7.6.3	Using Linear Transformations to Simplify Hypothesis Tests	83
7.6.4	Goodness of Fit	84
7.6.5	Using R^2 to test exclusion restrictions	86
7.7	Asymptotic Properties in Regression Models	88
7.7.1	Consistency of OLS Estimator	88
7.7.2	Linear Projection Model Revisited	90
7.7.3	Asymptotic Normality	91
7.7.4	Asymptotic Inference	95
8	Heteroskedasticity and Generalised Least Squares	100
8.1	Conditional Heteroskedasticity	100
8.2	Generalised Least Squares	102
8.2.1	Feasible Generalized Least Squares	105
8.2.2	Weighted Least Squares	106
8.3	Clustering	108
9	Endogeneity	111
9.1	Endogenous Explanatory Variable	111
9.2	Measurement Error	111
9.2.1	Multiple regression with errors in variables	115
9.3	Omitted Variable	116
9.3.1	Single Linear Regression with Omitted Variables	116
9.3.2	Multiple Regression with Omitted Variables	117

10 Instrumental Variables	119
10.1 Two Stage Least Squares	119
10.1.1 Just-Identified System	125
10.1.2 L Instruments	126
10.2 Control Function	128
10.3 Indirect Least Squares	129
10.4 Consistency of 2SLS Estimator	130
10.5 Asymptotics	132
10.6 IV with Multiple Regressors	133
10.7 Testing	134
10.7.1 Testing for Informativeness of Instrumental Variables	135
10.7.2 Test for Validity of Instrumental Variables	135
10.7.3 Test for Endogeneity	136
10.8 Finite Sample Problems	136
10.8.1 Overfitting	137
10.8.2 Weak Instrumental Variables	137
10.8.3 Anderson-Rubin Test	137
11 Generalised Method of Moments	139
11.1 Motivating Examples	139
11.1.1 GMM Applied to Linear Regression	139
11.1.2 GMM Applied to Linear Regression with Instrumental Variables	142
11.2 GMM General Framework	147
11.2.1 GMM Applied to IV Model	148
11.3 Asymptotics	149
11.3.1 Consistency of GMM Estimator	149
11.3.2 Asymptotic Normality of GMM Estimator	151
11.4 Identification	155
11.4.1 Identification in Linear Model with Instrumental Variables	156
11.5 Asymptotic Efficiency	157
11.5.1 Efficiency in Linear IV Model	158
11.6 Test For Overidentifying Restrictions	159
11.7 Hypothesis Testing	161
11.7.1 Linear Testing	161
11.7.2 Nonlinear Testing	162
11.8 Criterion-Based Test	164
12 Time Series	166
12.1 Introduction	166
12.2 Cross Sectional Data Review	166
12.2.1 Ordinary Least Squares	166
12.2.2 Maximum Likelihood Estimation	166
12.2.3 Partial Least Squares	168
12.2.4 Correlation	170

12.3	Autocorrelation	171
12.4	Partial Autocorrelation	174
12.5	Linear Projection in Time Series	177
12.5.1	Linear Time Series Models	180
12.6	Autoregression	184
12.7	Autoregressive Model With Intercepts and Time Trends	188
12.7.1	Random Walks	192
12.8	Stationarity	194
12.9	Unconditional Stationary Autoregression	198
12.10	Asymptotic Theory for Time Series	203
12.11	Limit Theorems for Martingale Difference Sequences	205
12.12	Limit Theorems for AR(1) Model	207
12.13	Higher Order Autoregression	215
12.13.1	AR(2) Model	215
12.13.2	AR(p) Model	221
12.14	ADL Models	225
12.15	VAR Models	229
12.15.1	Granger Causality	235
12.16	Misspecification Testing in Time Series	237
12.17	Unit Root Testing	238
12.17.1	Deterministic Terms	241
12.18	Cointegration	244
12.18.1	Beveridge-Nelson Decomposition	244
12.18.2	Cointegration	247
12.18.3	Cointegration for VAR(2)	249
12.18.4	Cointegration for VAR(p)	251
12.18.5	Granger Representation Theorem	254
13	Panel Data	256
13.1	Two-Way Error Components Model	256
13.1.1	Time Dummies	257
13.2	Pooled OLS in Linear Panel Model	258
13.3	Fixed Effects Model	263
13.3.1	Within Groups Estimator	264
13.3.2	Least Squares Dummy Variables Estimator	267
13.3.3	First-Differenced OLS	268
13.3.4	Efficiency Comparison of Fixed Effects Estimators	271
13.3.5	Large T Asymptotics of Fixed Effects Estimators	272
13.4	Random Effects Model	274
13.4.1	Between Groups Estimator	274
13.4.2	Prereqs For Random Effects Generalised Least Squares Estimator	276
13.4.3	Random Effects Generalised Least Squares Estimator	278
13.4.4	Theta Differencing	282
13.5	Testing for Random Effects	285

13.6 Time-Invariant Explanatory Variables	287
13.7 Models with Predetermined Covariates	288
13.7.1 First Differenced 2SLS Estimator	288

1 Probability Theory

1.1 Axioms

We define Ω to be the **sample space** containing all possible outcomes. However, we cannot simply assign probabilities to every subset of Ω ; we first specify a collection of measurable events. That is where a σ – *algebra* comes in.

Definition 1 (σ – algebra)

A collection of sets $\mathcal{A}$ is a σ – *algebra* if it satisfies the following properties:

1. $\Omega \in \mathcal{A}$;
2. If $A \in \mathcal{A}$ then the complement $A^c \in \mathcal{A}$;
3. If $A_1, A_2, \dots \in \mathcal{A}$, then

$$\bigcup_{n=1}^{\infty} A_n \in \mathcal{A}.$$

Now with a σ – *algebra*, we can define probability measures.

Definition 2 (Probability Measure)

A **probability measure** $\mathbb{P} : \mathcal{A} \rightarrow [0, 1]$ is a function satisfying:

1. $\mathbb{P}(\Omega) = 1$.
2. $\mathbb{P}(A) \in [0, 1]$ for all $A \in \mathcal{A}$.
3. If $A_1, A_2, \dots \in \mathcal{A}$ are disjoint events, then

$$\mathbb{P}\left(\bigcup_{i=1}^{\infty} A_i\right) = \sum_{i=1}^{\infty} \mathbb{P}(A_i).$$

A triplet $(\Omega, \mathcal{A}, \mathbb{P})$ is called a **probability space**.

Example 3

(Real Line). Let $\Omega = \mathbb{R}$. Then, the σ – *algebra* is given by the **Borel field**, which is the σ – *algebra* generated by intervals of the form $(-\infty, x]$. A probability measure over the Borel field induces the **cumulative distribution function** $F(x) = \mathbb{P}((-\infty, x])$.

We can now introduce random variables.

Definition 4 (Random Variable)

A random variable X is a measurable mapping $X : \Omega \rightarrow \mathbb{R}$ such that, for every Borel set $B \in \mathcal{B}$,

$$X^{-1}(B) \in \mathcal{A}.$$

1.2 Distribution Functions and Densities

We let X be a **random variable**. Every random variable has a distribution, which can be described by its **cumulative distribution function** $F(x) = \mathbb{P}(X \leq x)$. A discrete random variable has a **probability mass function**; when the distribution is absolutely continuous, it has a **probability density function** $f(x)$. In the latter case,

$$F(x) = \int_{-\infty}^x f(t) dt$$

Definition 5 (Cumulative Distribution Function)

A function $F(x)$ for $x \in \mathbb{R}$ is a **cumulative distribution function** if it satisfies:

1. $\lim_{x \rightarrow -\infty} F(x) = 0$ and $\lim_{x \rightarrow \infty} F(x) = 1$;
2. F is non-decreasing: $F(x) \leq F(y)$ for $x \leq y$;
3. $F(x)$ is right-continuous: $\lim_{h \downarrow 0} F(x + h) = F(x)$.

The cumulative distribution function uniquely describes the probability distribution of a random variable X . One way to think about this is that a distribution is the blueprint for a house whilst a random variable is a particular instance of a house.

A distribution can have many random variables associated with it but a random variable can only have one distribution associated to it.

Random Variable $\xrightarrow{\text{Induces}}$ Distribution $\xrightarrow{\text{Described by}}$ CDF (and, when it exists, a PMF or PDF)

Example 6

(Normal Random Variable). When we say that X is a normal random variable, what we are actually saying is that the random variable X has a **probability density function** of the form

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}$$

and the associated cumulative distribution function

$$F(x) = \int_{-\infty}^x f(t) dt = \int_{-\infty}^x \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{t-\mu}{\sigma}\right)^2} dt$$

Definition 7 (Atom)

An **atom** is defined as:

$$\mathbb{P}(X = x) = \mathbb{P}(X \leq x) - \mathbb{P}(X < x) = F(x) - \lim_{h \downarrow 0} F(x - h).$$

Definition 8

(Random Sample). The random variables $X_1, \dots, X_n$ are called a **random sample** of size n from a population if they are independent and identically distributed. If their common distribution has a density, then $f_i(x) = f(x)$ for $i = 1, \dots, n$.

Definition 9 (Joint Distribution Function)

A function $F(x_1, x_2) = \mathbb{P}(X_1 \leq x_1, X_2 \leq x_2)$ for $x_1, x_2 \in \mathbb{R}$ is a **joint cumulative distribution function** if:

1. $F(x_1, x_2) \rightarrow 0$ if either coordinate tends to $-\infty$, and $F(x_1, x_2) \rightarrow 1$ as both coordinates tend to ∞ ;
2. $F(x_1, x_2)$ is non-decreasing in x_1, x_2 ;
3. $F(x_1, x_2)$ is right continuous in x_1, x_2 ;
4. the probability assigned to every rectangle $(a_1, b_1] \times (a_2, b_2]$ is non-negative.

The **marginal distribution function** is defined as:

$$F_1(x_1) = \lim_{x_2 \rightarrow \infty} F(x_1, x_2). \quad (1)$$

Definition 10 (Independence)

Random variables $X_1, \dots, X_p$ are **independent** if their joint distribution function satisfies:

$$F(x_1, \dots, x_p) = \prod_{i=1}^p F_i(x_i) \quad (2)$$

for all $x_1, \dots, x_p$.

Proposition 11 (Joint Density Function of Random Sample)

With a random sample $(X_1, \dots, X_n)$, we can use the independent and identically distributed assumption to write the **joint density** of our random sample as:

$$f(x_1, \dots, x_n) = f_1(x_1)f_2(x_2)\dots f_n(x_n) = \prod_{i=1}^n f(x_i) \quad (3)$$

Proof. The first equality follows from independence and the second equality follows from being identically distributed. $\square$

Generally, when we have a random variable X with a probability density function $f(x)$ which is parameterised by θ , we usually write the probability density function as $f(x; \theta)$.

Example 12

(Joint density of exponential distribution). The density function of an exponential distribution is

$$f(x; \beta) = \frac{1}{\beta} e^{-x/\beta}, \quad x \geq 0, \beta > 0.$$

Then, the joint distribution is

$$\begin{aligned} f(x_1, \dots, x_n; \beta) &= \prod_{i=1}^n f(x_i; \beta) \\ &= \prod_{i=1}^n \frac{1}{\beta} e^{-x_i/\beta} = \frac{1}{\beta^n} e^{-(x_1 + \dots + x_n)/\beta} \end{aligned}$$

Now, suppose that we have our random sample $X_1, \dots, X_n$, we are interested in **functions** of our random sample.

Definition 13 (Statistic)

Let $X_1, \dots, X_n$ be a random sample of size n from a population. Let $T(x_1, \dots, x_n)$ be a measurable real-valued or vector-valued **function** defined on the sample space. Then,

$$Y = T(X_1, \dots, X_n)$$

is called a **statistic**. The probability distribution of a statistic Y is called the **sampling distribution** of Y .

Remark 14. Since $X_1, \dots, X_n$ are random variables, the statistic $Y = T(X_1, \dots, X_n)$ is also a **random variable**. That is why the statistic has a distribution (and a density when one exists).

Example 15

(Sample Mean) The sample mean is the classic example of a statistic

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$$

Example 16

(Sample Variance) The sample variance is given by

$$S^2 = \frac{1}{(n-1)} \sum_{i=1}^n (X_i - \bar{X})^2 \quad (4)$$

Definition 17 (Conditional Density)

For $f_X(x) > 0$, the conditional density of Y given $X = x$ is:

$$f_{Y|X}(y|x) = \frac{f_{X,Y}(x, y)}{f_X(x)} \quad (5)$$

We can rewrite this as:

$$f_{X,Y}(x, y) = f_{Y|X}(y|x)f_X(x)$$

If x and y are independent, then:

$$f_{Y|X}(y|x) = f_Y(y)$$

since $f_{X,Y}(x,y) = f_X(x)f_Y(y)$.

Definition 18 (Bayes Formula)

Bayes formula is given by:

$$f_{X|Y}(x|y) = \frac{f_{Y|X}(y|x)f_X(x)}{f_Y(y)} = \frac{f_{Y|X}(y|x)f_X(x)}{\int_{\mathbb{R}} f_{Y|X}(y|t)f_X(t)dt} \quad (6)$$

We now introduce some important probability density functions and cumulative distribution functions.

2 Transformations and Expectations

2.1 Transformations of Distributions

We state a useful theorem that we will use quite a few times.

Theorem 19 (Independence under Transformations)

Let $X_1, \dots, X_n$ be independent random variables. Suppose that $g_1, \dots, g_n$ are measurable functions from $\mathbb{R}$ to $\mathbb{R}$ and define $Y_i = g_i(X_i)$ for $i = 1, \dots, n$. Then, the random variables $Y_1, \dots, Y_n$ are independent.

Proof. Measurable functions of independent random variables are independent. $\square$

Suppose we have a random variable X with a density f_X . Furthermore, suppose that we have a monotone function $g(\cdot)$ such that we construct the new random variable $Y = g(X)$. What is the density function of Y ? The following theorem tells us this answer.

Theorem 20 (Change of Variables Formula)

Let X be a continuous random variable with density f_X and support $\mathcal{X}$. Let $g(\cdot)$ be a monotone mapping with a continuous derivative. Define the random variable $Y = g(X)$ with support $\mathcal{Y} = g(\mathcal{X})$. Then, the density of the random variable Y is given by:

$$f_Y(y) = f_X\left(g^{-1}(y)\right) \left| \frac{\partial}{\partial y} g^{-1}(y) \right| \quad (7)$$

for all $y \in \mathcal{Y}$.

Proof. If $g(\cdot)$ is strictly monotone, then its inverse exists. For an increasing transformation,

$$\frac{\partial}{\partial y}(g^{-1}(y)) = \frac{1}{g'(g^{-1}(y))}$$

We then have that:

$$\begin{aligned} F_Y(y) &= \mathbb{P}(Y \leq y) \\ &= \mathbb{P}(g(X) \leq y) \\ &= \mathbb{P}(X \leq g^{-1}(y)) \\ &= F_X(g^{-1}(y)) \end{aligned}$$

So we have:

$$F_Y(y) = F_X(g^{-1}(y))$$

Take derivatives of both sides with respect to y . Use the chain rule and plug in our expression for the derivative of the inverse:

$$f_Y(y) = \frac{d}{dy} \left[F_X(g^{-1}(y)) \right]$$

and hence

$$f_Y(y) = f_X(g^{-1}(y)) \left| \frac{1}{g'(g^{-1}(y))} \right|.$$

For a decreasing transformation the inequality is reversed, and the absolute value gives the same formula. $\square$

2.2 Expectations of Distributions

This section will essentially state a bunch of results we will need throughout the course.

Definition 21 (Expectation)

If $\mathbb{E}[|X|] < \infty$ then the expectation of the random variable X is:

$$\mathbb{E}[X] = \int_{\mathbb{R}} xf(x)dx.$$

Remark 22. *The Cauchy distribution has no expectation.*

Proposition 23 (Linearity)

The expectation operator is a linear operator:

$$\mathbb{E}[aX + b] = a\mathbb{E}[X] + b.$$

Definition 24 (Moments)

The r -th **moment** of a random variable X is given by:

$$\mathbb{E}[X^r]$$

for $r > 0$. The r -th **central moment** is given by:

$$\mathbb{E}[(X - \mathbb{E}[X])^r]$$

for $r > 0$.

Definition 25 (Variance)

The variance is given by:

$$\text{Var}[X] = \mathbb{E}[(X - \mathbb{E}[X])^2] = \mathbb{E}[X^2] - \mathbb{E}[X]^2 \quad (8)$$

Standard deviation is the square root of variance:

$$\text{SDV}(X) = \sqrt{\text{Var}[X]}.$$

Proposition 26 (Quadratic Variance)

The variance is a quadratic function:

$$\text{Var}[\mu + \sigma X] = \sigma^2 \text{Var}[X].$$

Definition 27 (Jensen's Inequality)

If g is convex, then:

$$g(\mathbb{E}[X]) \leq \mathbb{E}[g(X)].$$

Proposition 28 (Linear Expectation)

Let X be a random vector. Let A and B be matrices. The expectation is a linear function:

$$\mathbb{E}[AX + B] = A\mathbb{E}[X] + B.$$

Proposition 29 (Quadratic Variance)

Let X be a random vector. Let A and B be matrices. The variance is a quadratic function:

$$\text{Var}[AX + B] = A\text{Var}[X]A^T$$

Definition 30 (Covariance)

The covariance between X and Y is:

$$\text{Cov}[X, Y] = \mathbb{E}[XY] - \mathbb{E}[X]\mathbb{E}[Y]. \quad (9)$$

Proposition 31 (Covariance Properties)

Let X, Y be random variables. Let a and b be scalars. Then:

$$\text{Cov}[aX + b, cY + d] = ac \cdot \text{Cov}[X, Y]$$

If X and Y are independent, then

$$\text{Cov}[X, Y] = 0.$$

The converse is **not** true.

Proposition 32 (Variance of Two variables)

Let X, Y be random variables. Let a and b be scalars. Then:

$$\text{Var}[aX + bY] = a^2\text{Var}[X] + b^2\text{Var}[Y] + 2ab\text{Cov}(X, Y).$$

Definition 33 (Correlation)

Correlation is a normalised covariance:

$$\text{Corr}[X, Y] = \frac{\text{Cov}(X, Y)}{\sqrt{\text{Var}(X)\text{Var}(Y)}}$$

Definition 34 (Cauchy-Schwarz Inequality)

The Cauchy-Schwarz inequality is:

$$|\mathbb{E}[XY]|^2 \leq \mathbb{E}[X^2]\mathbb{E}[Y^2].$$

Definition 35 (Conditional Expectation)

Conditional expectation of Y given X is:

$$\mathbb{E}[Y|X = x] = \int_{-\infty}^{\infty} y f(y|x) dy.$$

Definition 36 (Law of Iterated Expectations)

$$\mathbb{E}[Y] = \mathbb{E}[\mathbb{E}[Y|X]]$$

$$\text{Var}[Y] = \mathbb{E}[\text{Var}[Y|X]] + \text{Var}[\mathbb{E}[Y|X]]$$

$$\text{Cov}(X, Y) = \mathbb{E}[\text{Cov}(X, Y|Z)] + \text{Cov}(\mathbb{E}[X|Z], \mathbb{E}[Y|Z])$$

2.3 Multivariate Normal Distributions

We now wish to define the multivariate normal distribution. This is a distribution we will use a lot throughout the course. First, we can define the standard multivariate normal distribution.

Definition 37 (Standard Multivariate Normal Distribution)

The random vector $\mathbf{X} = (X_1, \dots, X_p)^T$ is a standard multivariate normal distribution $\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$ if it has the density function:

$$f(\mathbf{x}) = (2\pi)^{-p/2} \exp(-\mathbf{x}^T \mathbf{x} / 2). \quad (10)$$

With the standard multivariate normal distribution, we can use the properties of the expectation and variance operators to define the regular multivariate normal distribution.

Definition 38 (Regular Multivariate Normal Distribution)

Let $\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$. Define the random variable $\mathbf{Y} = \boldsymbol{\mu} + \mathbf{A}\mathbf{X}$ for $\boldsymbol{\mu} \in \mathbb{R}^p$ and $\mathbf{A} \in \mathbb{R}^{p \times p}$. Denote the covariance matrix $\boldsymbol{\Omega} = \mathbf{A}\mathbf{A}^T$ and assume that it is invertible.

Then, $\mathbf{Y} \sim \mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Omega})$ with the density function:

$$f(\mathbf{y}) = \det(2\pi\boldsymbol{\Omega})^{-1/2} \exp \left\{ -\frac{1}{2}(\mathbf{y} - \boldsymbol{\mu})^T \boldsymbol{\Omega}^{-1}(\mathbf{y} - \boldsymbol{\mu}) \right\} \quad (11)$$

Notice that in the definition of a regular multivariate normal distribution, we required that the inverse of the covariance matrix $\boldsymbol{\Omega}^{-1}$ to exist. However, what happens if this inverse does not exist? Therefore, we can define what is known as a **singular** multivariate normal distribution to handle this.

Definition 39 (Singular Multivariate Normal Distribution)

Let Ω be a symmetric positive semidefinite covariance matrix. We say that $\mathbf{Y}$ follows a **singular normal distribution** $\mathbf{Y} \sim \mathcal{N}(\boldsymbol{\mu}, \Omega)$ if, for every vector $\mathbf{V} \in \mathbb{R}^p$, one of the following holds for $\mathbf{V}^T \mathbf{Y}$:

1. $\mathbf{V}^T \cdot \mathbf{Y}$ has a regular univariate normal density function:

$$\mathbf{V}^T \cdot \mathbf{Y} \sim \mathcal{N}(\mathbf{V}^T \boldsymbol{\mu}, \mathbf{V}^T \Omega \mathbf{V}) \quad (12)$$

- 2.

$$\mathbf{V}^T \cdot \mathbf{Y} = \mathbf{V}^T \cdot \boldsymbol{\mu} \quad (13)$$

This definition also applies when the covariance matrix Ω is not invertible.

The multivariate normal has a few nice properties.

Proposition 40 (Expectation and Variance)

Let $\mathbf{Y} \sim \mathcal{N}(\boldsymbol{\mu}, \Omega)$. Then:

$$\mathbb{E}[\mathbf{Y}] = \boldsymbol{\mu}$$

$$\text{Var}[\mathbf{Y}] = \Omega$$

Proposition 41 (Linear Transformations)

Let $\mathbf{Y} \sim \mathcal{N}(\boldsymbol{\mu}, \Omega)$. Then:

$$\mathbf{B} + \mathbf{C}\mathbf{Y} \sim \mathcal{N}(\mathbf{B} + \mathbf{C}\boldsymbol{\mu}, \mathbf{C}\Omega\mathbf{C}^T).$$

Proposition 42 (Independence and Zero Covariance for Jointly Normal Variables)

Let $(Y_1, Y_2)^T$ be jointly normally distributed. Then, Y_1 and Y_2 are independent if and only if $\text{Cov}(Y_1, Y_2) = 0$.

Proposition 43 (Convolution)

If $Y_1 \sim \mathcal{N}(\mu_1, \sigma_1^2)$ and $Y_2 \sim \mathcal{N}(\mu_2, \sigma_2^2)$ are independent, then:

$$Y_1 + Y_2 \sim \mathcal{N}(\mu_1 + \mu_2, \sigma_1^2 + \sigma_2^2).$$

2.4 Normal Sampling Distributions

2.4.1 Chi-Squared Distribution

We now move onto one of the most useful distributions besides the normal distribution in this course.

Definition 44 (Chi-squared Random Variable)

A random variable χ_p^2 is a chi-squared random variable with p degrees of freedom if it has the probability density function

$$f(x) = \frac{1}{\Gamma(p/2)2^{p/2}} x^{(p/2)-1} e^{-x/2}$$

for $x \in (0, \infty)$. The p here is the **degrees of freedom** of the χ^2 distribution.

Remark 45. We never really work with the density function of the χ^2 random variable in practice, but do know that it has one!

We come to a **very important** result we will need throughout the course.

Theorem 46 (Multivariate Normal and Chi-Squared Distribution)

Suppose that we have a random variable $\mathbf{X}$ which has a multivariate normal distribution

$$\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \Sigma)$$

whereby $\mathbf{X}$ has dimension d and the covariance matrix Σ has full rank d .

Then, we have that the following quadratic form of $\mathbf{X}$ has the distribution:

$$\mathbf{X}^T \Sigma^{-1} \mathbf{X} \sim \chi^2(d) \quad (14)$$

which is a χ^2 distribution with d degrees of freedom.

Proof. Recall from linear algebra.

Theorem 47 (Eigendecomposition of a matrix)

Let Σ be a real symmetric matrix. Then, we can express this matrix in terms of its eigendecomposition:

$$\Sigma = \mathbf{V} \mathbf{D} \mathbf{V}^T \quad (15)$$

where $\mathbf{V}$ is an orthogonal matrix containing eigenvectors in every column and $\mathbf{D}$ is a diagonal matrix containing the corresponding eigenvalues.

We also need another little fact about orthogonal matrices.

Lemma 48

If $\mathbf{V}$ is an orthogonal matrix, then

$$\mathbf{V} \mathbf{V}^T = \mathbb{I}$$

whereby the transpose of $\mathbf{V}$ is the inverse of $\mathbf{V}$.

By our result on eigendecomposition, we can write the inverse of the covariance matrix as:

$$\Sigma^{-1/2} = \mathbf{V} \mathbf{D}^{-1/2} \mathbf{V}^T$$

whereby $\mathbf{D}^{-1/2}$ is simply taking the inverse square root of the diagonal entries of $\mathbf{D}$.

We now apply a change of variables whereby for $\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \Sigma)$, we can write the rescaled random variable:

$$\mathbf{Z} = \Sigma^{-1/2} \mathbf{X}$$

We now want to derive the expectation and variance of this new random variable $\mathbf{Z}$.

$$\mathbb{E}[\mathbf{Z}] = \mathbb{E}[\Sigma^{-1/2} \mathbf{X}] = \Sigma^{-1/2} \mathbb{E}[\mathbf{X}] = \mathbf{0}$$

Since the mean of $\mathbf{Z}$ is zero, its covariance matrix is its second-moment matrix:

$$\begin{aligned}\mathbb{E}[\mathbf{Z}\mathbf{Z}^T] &= \mathbf{\Sigma}^{-1/2}\mathbb{E}[\mathbf{X}\mathbf{X}^T]\mathbf{\Sigma}^{-1/2} \\ &= \mathbf{\Sigma}^{-1/2}\text{Var}[\mathbf{X}]\mathbf{\Sigma}^{-1/2} \\ &= \mathbf{\Sigma}^{-1/2}\mathbf{\Sigma}\mathbf{\Sigma}^{-1/2} \\ &= \mathbf{V}\mathbf{D}^{-1/2}\mathbf{V}^T\mathbf{V}\mathbf{D}\mathbf{V}^T\mathbf{V}\mathbf{D}^{-1/2}\mathbf{V}^T \\ &= \mathbf{V}\mathbf{D}^{-1/2}\mathbf{D}\mathbf{D}^{-1/2}\mathbf{V}^T = \mathbf{V}\mathbf{V}^T = \mathbb{I}\end{aligned}$$

We therefore have that $\mathbf{Z} \sim \mathcal{N}(\mathbf{0}, \mathbb{I})$. Therefore, the sum of squares of $\mathbf{Z}$ is a χ^2 -distribution. We can therefore conclude that:

$$\begin{aligned}\mathbf{Z}^T\mathbf{Z} &= \mathbf{X}^T\mathbf{\Sigma}^{-1/2}\mathbf{\Sigma}^{-1/2}\mathbf{X} \\ &= \mathbf{X}^T\mathbf{\Sigma}^{-1}\mathbf{X} \sim \chi^2(d).\end{aligned}$$

□

We can extend this result to when the mean $\boldsymbol{\mu} \neq \mathbf{0}$.

Corollary 49

Suppose that we have a random variable $\mathbf{X}$ which has a multivariate normal distribution

$$\mathbf{X} \sim \mathcal{N}(\boldsymbol{\mu}, \mathbf{\Sigma})$$

whereby $\mathbf{X}$ has dimension d and the covariance matrix $\mathbf{\Sigma}$ has full rank d .

Then, we have that the following quadratic form of $\mathbf{X}$ has the distribution:

$$[\mathbf{X} - \boldsymbol{\mu}]^T \mathbf{\Sigma}^{-1} [\mathbf{X} - \boldsymbol{\mu}] \sim \chi^2(d) \quad (16)$$

which is a χ^2 distribution with d degrees of freedom.

Proposition 50 (Quadratic Form and Chi-squared Distribution)

If the $n \times 1$ vector $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ and the non-stochastic $n \times n$ matrix $\mathbf{G}$ is symmetric and idempotent with $\text{rank}(\mathbf{G}) = r \leq n$, then the scalar:

$$w = \mathbf{z}^T \mathbf{G} \mathbf{z} \sim \chi^2(r). \quad (17)$$

The following result is a useful relationship between standard normal random variables and χ^2 random variables.

Proposition 51 (Chi-squared is the sum of squared standard normals)

Let $Z_i \sim \mathcal{N}(0, 1)$ for $i = 1, \dots, n$ be independent standard normal random variables. Then

$$\mathbf{Z}^T \mathbf{Z} = \sum_{i=1}^n Z_i^2 \sim \chi^2(n)$$

That is, the sum of n squared standard normal random variables is a χ^2 random variable with n degrees of freedom.

We get the useful corollary.

Proposition 52 (Chi-squared is the square of a standard normal)

If $Z \sim \mathcal{N}(0, 1)$ then

$$Z^2 \sim \chi^2(1)$$

That is, the square of a standard normal random variable is a χ^2 random variable with 1 degree of freedom.

Proposition 53 (Chi-Squared and Orthogonal Projections)

Let $\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)$ and let $\mathbf{P}$ and $\mathbf{Q}$ be symmetric, idempotent matrices. Then,

$$\mathbf{X}^T \mathbf{P} \mathbf{X} \sim \chi^2(\text{tr}(\mathbf{P}))$$

$$\mathbf{X}^T \mathbf{Q} \mathbf{X} \sim \chi^2(\text{tr}(\mathbf{Q}))$$

Furthermore, if $\mathbf{PQ} = \mathbf{0}$, then $\mathbf{X}^T \mathbf{P} \mathbf{X}$ and $\mathbf{X}^T \mathbf{Q} \mathbf{X}$ are independent.

2.4.2 T and F Distributions**Definition 54** (Student's t distribution)

The random variable T has a Student's t distribution with p degrees of freedom if T has the probability density function

$$f_T(t) = \frac{\Gamma(\frac{p+1}{2})}{\Gamma(\frac{p}{2})} \frac{1}{(p\pi)^{1/2}} \frac{1}{(1 + t^2/p)^{(p+1)/2}}$$

for $t \in (-\infty, \infty)$.

We have the following result that relates a standard normal random variable, χ^2 random variable, and t-distributed random variable.

Proposition 55 (Deriving a t-distribution from a normal and chi-squared distribution)

Let $X \sim \mathcal{N}(0, 1)$ be a standard normal random variable and let $Y \sim \chi^2(p)$ be independent of X . Then,

$$\frac{X}{\sqrt{Y/p}} \sim t_{(p)} \quad (18)$$

which is a t distribution with p degrees of freedom. The ratio of a standard normal random variable to the square root of an independent chi-squared random variable divided by its degrees of freedom has a Student t-distribution with those degrees of freedom.

Remark 56. This says that the probability density function of the ratio $\frac{X}{\sqrt{Y/p}}$ is the Student's t probability density function. Note that we need to divide the χ^2 random variable by its degrees of freedom for this result to hold.

Example 57

We know that if $X_1, \dots, X_n$ are from a $\mathcal{N}(\mu, \sigma^2)$ sample, then we have that

$$\frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \sim \mathcal{N}(0, 1) \quad (19)$$

In practice, we do not know σ . Let $\hat{\sigma}^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$ be the sample variance, which we can plug in:

$$\frac{\bar{X} - \mu}{\hat{\sigma}/\sqrt{n}} \quad (20)$$

The question we want to ask is what is the distribution of this new ratio. We can multiply and divide by σ to get:

$$\frac{(\bar{X} - \mu)/(\sigma/\sqrt{n})}{\sqrt{\hat{\sigma}^2/\sigma^2}}$$

Now, we know that the numerator is:

$$\frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \sim \mathcal{N}(0, 1)$$

and the rescaled sample variance satisfies

$$\frac{(n-1)\hat{\sigma}^2}{\sigma^2} \sim \chi^2(n-1).$$

Equivalently, $\sqrt{\hat{\sigma}^2/\sigma^2}$ is the square root of this chi-squared variable divided by $n-1$. It is independent of the numerator, so the Student-t result implies that:

$$\frac{\bar{X} - \mu}{\hat{\sigma}/\sqrt{n}} \sim t(n-1).$$

Definition 58 (F-distribution)

If the random variables $w_1 \sim \chi^2(m)$ and $w_2 \sim \chi^2(n)$ and w_1 and w_2 are independent, then the scalar:

$$\nu = \frac{(w_1/m)}{(w_2/n)} \sim F(m, n). \quad (21)$$

The ratio of two independent chi-squared random variables, each divided by their respective degrees of freedom, has a F-distribution with the degrees of freedom of the numerator and denominator respectively.

3 Linear Regression Model and Linear Projection Model

3.1 Conditional Expectation Function

We give a motivation for why we care so much about conditional expectations and how linear projection relates to them. Generally, we have a dependent variable of interest Y and want to learn about the factors driving it. For example, Y could be an individual's wage and X could be an individual's sex. We may then ask what the average wage of a female employee is and compare it with the average wage of a male employee. Using mathematics, we can express this using conditional expectations:

$$\mathbb{E}[\text{wage}|\text{sex} = \text{female}]$$

$$\mathbb{E}[\text{wage}|\text{sex} = \text{male}]$$

To generalise this, we are interested in the conditional expectation of Y given values that the variables $X_1, \dots, X_k$ can take on:

$$\mathbb{E}[Y|X_1 = x_1, X_2 = x_2, \dots, X_k = x_k] = m(x_1, \dots, x_k).$$

We can write the vector $\mathbf{X} = \begin{bmatrix} X_1 \\ X_2 \\ \vdots \\ X_k \end{bmatrix} \in \mathbb{R}^k$ and from this, write the conditional expectation as:

$$\mathbb{E}[Y|\mathbf{X} = \mathbf{x}] = m(\mathbf{x}).$$

Definition 59 (Conditional Expectation Function)

The conditional expectation function $m(\mathbf{x})$ is a function of $\mathbf{x} = (x_1, \dots, x_k)$:

$$m(\mathbf{x}) = \mathbb{E}[Y|\mathbf{X} = \mathbf{x}]. \quad (22)$$

Remark 60. We can interpret this as when $\mathbf{X}$ takes on the value $\mathbf{x}$, then the average value of Y is $m(\mathbf{x})$.

We can also write the random variable:

$$m(\mathbf{X}) = \mathbb{E}[Y|\mathbf{X}].$$

Now, given the joint density $f(y, \mathbf{x})$, recall that the conditional density is given by:

$$f_{Y|\mathbf{X}}(y|\mathbf{x}) = \frac{f(y, \mathbf{x})}{f_{\mathbf{X}}(\mathbf{x})}.$$

Then, this means we can write the CEF of Y given $\mathbf{X} = \mathbf{x}$ as:

$$m(\mathbf{x}) = \mathbb{E}[Y|\mathbf{X} = \mathbf{x}] = \int_{-\infty}^{\infty} y \cdot f_{Y|\mathbf{X}}(y|\mathbf{x}) dy.$$

which looks at all possible values of y given the probability of that value y , $f_{Y|\mathbf{X}}(y|\mathbf{x})$. We can interpret $m(\mathbf{x})$ as the

expectation of Y for a subpopulation whereby $\mathbf{X} = \mathbf{x}$. We can plug in $\mathbf{X}$ to get:

$$m(\mathbf{X}) = \mathbb{E}[Y|\mathbf{X}] = \int_{-\infty}^{\infty} y \cdot f_{Y|\mathbf{X}}(y|\mathbf{X}) dy.$$

To recap, we have the variable Y and $\mathbf{X}$. From this, we constructed the CEF function $m(\mathbf{X})$. We can then define the CEF error.

Definition 61 (CEF Error)

The conditional expectation function error e is defined as the difference between Y and the CEF:

$$e \equiv Y - m(\mathbf{X}). \quad (23)$$

From this, we get the formula:

$$Y = m(\mathbf{X}) + e. \quad (24)$$

Proposition 62 (Properties of CEF Error)

Let e be the CEF error. Then,

1. $\mathbb{E}[e|\mathbf{X}] = 0$
2. $\mathbb{E}[e] = 0$
3. For any function $h(\mathbf{X})$, $\mathbb{E}[h(\mathbf{X})e] = 0$.

Proof. To show the first result:

$$\begin{aligned} \mathbb{E}[e|\mathbf{X}] &= \mathbb{E}[(Y - m(\mathbf{X}))|\mathbf{X}] \\ &= \mathbb{E}[Y|\mathbf{X}] - \mathbb{E}[m(\mathbf{X})|\mathbf{X}] \\ &= m(\mathbf{X}) - m(\mathbf{X}) = 0. \end{aligned}$$

To show the second result, we use the law of iterated expectations:

$$\mathbb{E}[e] = \mathbb{E}[\mathbb{E}[e|\mathbf{X}]] = \mathbb{E}[0] = 0.$$

For the last one:

$$\mathbb{E}[h(\mathbf{X})e] = \mathbb{E}[\mathbb{E}[h(\mathbf{X})e|\mathbf{X}]] = \mathbb{E}[h(\mathbf{X})\mathbb{E}[e|\mathbf{X}]] = \mathbb{E}[h(\mathbf{X}) \cdot 0] = 0.$$

□

This third property has a well-known name.

Definition 63 (Orthogonality Condition)

Let e be the CEF error. Then the condition:

$$\mathbb{E}[h(\mathbf{X})e] = 0 \quad (25)$$

for any measurable function $h(\cdot)$ is called the **orthogonality condition**.

The orthogonality condition implies that the covariance of $\mathbf{X}$ and e is zero.

Proposition 64 (Covariance of Regressor with CEF Error is Zero)

Suppose that the orthogonality condition holds $\mathbb{E}[h(\mathbf{X})e] = 0$. Then, the covariance of the regressor $\mathbf{X}$ with the CEF error is zero:

$$\text{Cov}[\mathbf{X}, e] = 0. \quad (26)$$

Proof. The covariance is:

$$\text{Cov}[\mathbf{X}, e] = \mathbb{E}[\mathbf{X}e] - \mathbb{E}[\mathbf{X}]\mathbb{E}[e] = 0 - 0 = 0$$

as $\mathbb{E}[\mathbf{X}e] = \mathbb{E}[h(\mathbf{X})e] = 0$ by the orthogonality condition where $h(\mathbf{X}) = \mathbf{X}$. □

We have a special name for the variance of the error term e .

Definition 65 (Unconditional Variance of Error Term)

The **unconditional variance of the error term** is:

$$\sigma^2 = \text{Var}[e]. \quad (27)$$

Proof. We have that:

$$\begin{aligned} \sigma^2 &= \text{Var}[e] \\ &= \mathbb{E}[e^2] - \mathbb{E}[e]^2 \\ &= \mathbb{E}[e^2] - 0 \\ &= \mathbb{E}[e^2] \end{aligned}$$

where e^2 is the second moment of error. □

Remark 66. We can interpret σ^2 as the amount of variation in y that is not accounted for in the conditional expectation $\mathbb{E}[Y|\mathbf{X}]$.

Intuitively, as you include more variables in your conditional expectation function, then the variance of your CEF error should drop.

Now, suppose we have the variables y and $\mathbf{X}$, and we want to predict Y using a function $g(\mathbf{X})$. We want to know what is the function $g(\cdot)$ that we can use as our best predictor (under mean squared error). We state the optimal function for this task.

Proposition 67 (Conditional Expectation Minimises the MSE)

The conditional expectation $\mathbb{E}[Y|\mathbf{X}]$ minimises the mean-squared error:

$$\mathbb{E}[(Y - g(\mathbf{X}))^2] \quad (28)$$

where $g(\mathbf{X}) = \mathbb{E}[Y|\mathbf{X}]$.

Proof. We want to find a $g(\mathbf{X})$ that minimises the mean squared error. Expand out the square:

$$\begin{aligned}\mathbb{E}[(Y - g(\mathbf{X}))^2] &= \mathbb{E}[(m(\mathbf{X}) + e - g(\mathbf{X}))^2] \\ &= \mathbb{E}[e^2] + 2\mathbb{E}[e(m(\mathbf{X}) - g(\mathbf{X}))] + \mathbb{E}[(m(\mathbf{X}) - g(\mathbf{X}))^2] \\ &=_{(1)} \mathbb{E}[e^2] + 0 + \mathbb{E}[(m(\mathbf{X}) - g(\mathbf{X}))^2] \\ &= \mathbb{E}[e^2] + \mathbb{E}[(m(\mathbf{X}) - g(\mathbf{X}))^2]\end{aligned}$$

where $=_{(1)}$ uses the orthogonality condition. Now, if we look at the last equation:

$$\mathbb{E}[e^2] + \mathbb{E}[(m(\mathbf{X}) - g(\mathbf{X}))^2]$$

we can see that we cannot choose a $g(\mathbf{X})$ to alter the first term, only the second term. The most we can do to minimise this equation is to set:

$$g(\mathbf{X}) = m(\mathbf{X})$$

as this sets the second term to be 0 (cannot be negative due to the square). $\square$

Hence, we have shown that the conditional expectation function is the function that minimises the mean squared error. That is, the best predictor of Y is $E[Y|\mathbf{X}]$.

3.2 Linear Regression Model

We can now define a special model. Now, if we impose the assumption that the conditional expectation function **has** to be a linear function of $\mathbf{X}$, which we can denote by $\mathbf{X}^T\boldsymbol{\beta}$, then we get the linear regression model.

Definition 68 (Linear Regression Model)

The Linear Regression model is given by:

$$Y = \mathbf{X}^T\boldsymbol{\beta} + e \tag{29}$$

$$\mathbb{E}[e|\mathbf{X}] = 0. \tag{30}$$

Remark 69. *This linear conditional-mean specification is a special case of the general regression model discussed later.*

Under the linear regression model, $\mathbb{E}[e|\mathbf{X}] = 0$. We have already shown that this implies the error term is also unconditionally mean zero and orthogonal to the regressors:

$$\mathbb{E}[e] = 0$$

$$\mathbb{E}[\mathbf{X}e] = \mathbf{0}.$$

3.3 Linear Projection Model

We now want to derive a different model that is related to the linear regression model. We shall motivate it from another perspective. We know that $\mathbf{X}^T\boldsymbol{\beta}$ is a linear predictor of Y . The question we now want to ask is: what value of $\boldsymbol{\beta}$ minimises the mean squared error of predicting Y ? This value is known as the **linear projection coefficient**.

Proposition 70 (Linear Projection Coefficient)

Let $\mathbf{X}^T\boldsymbol{\beta}$ be our linear predictor for Y . The **linear projection coefficient** minimises the mean squared error of predicting Y :

$$\min_{\boldsymbol{\beta}} \mathbb{E}[(Y - \mathbf{X}^T\boldsymbol{\beta})^2]. \quad (31)$$

The expression for $\boldsymbol{\beta}$ that minimises this is:

$$\boldsymbol{\beta} = (\mathbb{E}[\mathbf{X}\mathbf{X}^T])^{-1}\mathbb{E}[\mathbf{X}Y]. \quad (32)$$

Proof. First, write out the mean squared error:

$$\mathbb{E}[(Y - \mathbf{X}^T\boldsymbol{\beta})^2] = \mathbb{E}[Y^2] - 2\boldsymbol{\beta}^T\mathbb{E}[\mathbf{X}Y] + \boldsymbol{\beta}^T\mathbb{E}[\mathbf{X}\mathbf{X}^T]\boldsymbol{\beta}.$$

Now, we take first order conditions of this. To do so, recall from matrix calculus:

$$\begin{aligned} \frac{\partial}{\partial \boldsymbol{\beta}} \boldsymbol{\beta}^T &= \mathbf{I} \\ \frac{\partial}{\partial \boldsymbol{\beta}} \boldsymbol{\beta}^T \mathbf{A} \boldsymbol{\beta} &= 2\mathbf{A}\boldsymbol{\beta}. \end{aligned}$$

So taking first order conditions, we have that:

$$2\mathbb{E}[\mathbf{X}\mathbf{X}^T]\boldsymbol{\beta} = 2\mathbb{E}[\mathbf{X}Y]$$

Now, assuming invertibility, we can conclude that:

$$\boldsymbol{\beta} = \mathbb{E}[\mathbf{X}\mathbf{X}^T]^{-1}\mathbb{E}[\mathbf{X}Y].$$

□

Hence, we have argued that the best linear predictor of Y is:

$$\mathbf{X}^T\boldsymbol{\beta} = \mathbf{X}^T\mathbb{E}[\mathbf{X}\mathbf{X}^T]^{-1}\mathbb{E}[\mathbf{X}Y]$$

We now notice that under this definition, we still have the orthogonality condition:

$$\begin{aligned} \mathbb{E}[\mathbf{X}e] &= \mathbb{E}[\mathbf{X}(Y - \mathbf{X}^T\boldsymbol{\beta})] \\ &= \mathbb{E}[\mathbf{X}Y - \mathbf{X}\mathbf{X}^T\boldsymbol{\beta}] \\ &= \mathbb{E}[\mathbf{X}Y] - \mathbb{E}[\mathbf{X}\mathbf{X}^T]\boldsymbol{\beta} \\ &= \mathbb{E}[\mathbf{X}Y] - \mathbb{E}[\mathbf{X}\mathbf{X}^T]\boldsymbol{\beta} \\ &= \mathbb{E}[\mathbf{X}Y] - \mathbb{E}[\mathbf{X}\mathbf{X}^T]\mathbb{E}[\mathbf{X}\mathbf{X}^T]^{-1}\mathbb{E}[\mathbf{X}Y] \\ &= \mathbb{E}[\mathbf{X}Y] - \mathbf{I}\mathbb{E}[\mathbf{X}Y] \\ &= \mathbf{0} \end{aligned}$$

Now, we can define our alternative model.

Definition 71 (Linear Projection Model)

The **linear projection model** is defined by:

$$\begin{aligned} Y &= \mathbf{X}^T \boldsymbol{\beta} + e \\ \mathbb{E}[\mathbf{X}e] &= \mathbf{0} \\ \boldsymbol{\beta} &= (\mathbb{E}[\mathbf{X}\mathbf{X}^T])^{-1} \mathbb{E}[\mathbf{X}Y]. \end{aligned}$$

When you have a linear projection model, you are always guaranteed by definition:

$$\mathbb{E}[\mathbf{X}e] = \mathbf{0}.$$

Therefore, we still have that result that every regressor x_j is orthogonal to the error term e . Now, if we include an intercept $x_j = 1$, then we get that:

$$\mathbb{E}[1 \cdot e] = \mathbb{E}[e] = 0.$$

Hence, we can see now why we include an intercept in a linear regression model. By including an intercept term, we are guaranteed that the error term has mean 0 by the orthogonality condition:

$$\mathbb{E}[e] = 0.$$

If we don't have this, we no longer have any guarantees that the error term is mean zero.

Now, you should notice that strict exogeneity $\mathbb{E}[e|\mathbf{X}] = 0$ implies the orthogonality condition $\mathbb{E}[\mathbf{X}e] = \mathbf{0}$. This gives us the useful result that relates the linear regression and linear projection models we have defined.

Proposition 72 (Linear Regression and Linear Projection Model)

The linear regression model implies the linear projection model.

Remark 73. *The converse is not necessarily true. The projection error may not necessarily satisfy:*

$$\mathbb{E}[e|\mathbf{X}] = 0.$$

Therefore, the linear projection model is a weaker model compared to the linear regression model.

4 Regression and Matrix Algebra

Suppose we have a sample of n outcome variables (such as wage) we are interested in analysing. We can stack these outcome variables into a $n \times 1$ vector, which we denote by $\mathbf{y}$ as:

$$\mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix}.$$

Suppose that we have K explanatory variables of interest and n observations that may explain the outcome variable $\mathbf{y}$. We can stack these variables into a $n \times K$ data matrix:

$$\mathbf{X} = \begin{bmatrix} x_{11} & \dots & x_{1K} \\ x_{21} & \dots & x_{2K} \\ \vdots & \dots & \vdots \\ x_{n1} & \dots & x_{nK} \end{bmatrix}$$

We will be working with this matrix a lot. We now state some definitions and properties of matrices that will be useful to us.

Definition 74 (Symmetric)

A matrix $\mathbf{A}$ is symmetric if:

$$\mathbf{A}^T = \mathbf{A}.$$

Proposition 75 (Symmetry of Matrices)

For any matrix $\mathbf{A}$, the matrices $\mathbf{A}^T \mathbf{A}$ and $\mathbf{A} \mathbf{A}^T$ are **symmetric matrices**.

Definition 76 (Idempotent)

A square matrix $\mathbf{A}$ is idempotent if:

$$\mathbf{A} \mathbf{A} = \mathbf{A}.$$

Proposition 77 (Symmetric and Idempotent Matrices)

If $\mathbf{A}$ is a symmetric and idempotent matrix, then

$$\mathbf{A}^T \mathbf{A} = \mathbf{A}.$$

Proof.

$$\mathbf{A}^T \mathbf{A} = \mathbf{A} \mathbf{A} = \mathbf{A}$$

where the first equality uses symmetry and the second equality uses idempotency. □

Going back to our data matrix $\mathbf{X}$, a key assumption we make throughout the course is that we have more observations than explanatory variables. This makes sense or else we will not have enough degrees of freedom to estimate the parameters on our K explanatory variables if we have more parameters than observations.

Assumption 78 (Number of Observations)

For the $n \times K$ data matrix $\mathbf{X}$, we have that $n > K$.

We are now interested in the question of when is the matrix $\mathbf{X}^T \mathbf{X}$ invertible (this will become clear when we perform OLS). To do so, we need to build up some more definitions.

Definition 79 (Linear Independence)

Let $\mathbf{X}$ be a $n \times K$ matrix. The **column rank** of $\mathbf{X}$ is the number of **linearly independent columns**. The **row rank** of $\mathbf{X}$ is the number of **linearly independent rows**.

We show an important result regarding the rank of a matrix.

Theorem 80 (Column Rank Equals Row Rank)

The row rank is equal to the column rank.

Hence, so far, we know that as $\mathbf{X}$ is a $n \times K$ matrix, we have that:

$$\text{rank}(\mathbf{X}) \leq \min\{n, K\}$$

However, we generally assume that we have more observations than regressors in linear regression $K < n$.

Therefore, this implies that the rank of $\mathbf{X}$ is:

$$\text{rank}(\mathbf{X}) \leq K.$$

Therefore, if $\text{rank}(\mathbf{X}) = K$, then $\mathbf{X}$ has **full rank**. In the case that $\text{rank}(\mathbf{X}) < K$, this has the interpretation that the model has **perfect multicollinearity**.

Proposition 81 (Rank of Design Matrix)

A useful result for the design matrix is:

$$\text{rank}(\mathbf{X}^T \mathbf{X}) = \text{rank}(\mathbf{X})$$

where $\mathbf{X}$ is a $n \times K$ matrix and $\mathbf{X}^T \mathbf{X}$ is a $K \times K$ matrix.

However, it may be hard to compute the rank so we can compute the **determinant** instead. We have a result that relates the rank of a matrix and the determinant of a matrix.

Theorem 82 (Rank and Determinant of a Matrix)

A square matrix $\mathbf{A}$ has full rank if and only if $\det(\mathbf{A}) \neq 0$.

Therefore, we can pull the three concepts of invertibility, rank, and determinant of a matrix $\mathbf{A}$ as follows:

1. $\mathbf{A}$ is invertible

2. $\mathbf{A}$ is full rank

3. $\det(\mathbf{A}) \neq 0$

This applies to our matrix $\mathbf{X}^T \mathbf{X}$.

4.1 Ordinary Least Squares

Let $\mathbf{u}$ be the $n \times 1$ **residual vector**. An important thing to note is the order of the vector multiplication.

If we write $\mathbf{u}^T \mathbf{u}$, this gives us the **residual sum of squares** as:

$$\mathbf{u}^T \mathbf{u} = \begin{bmatrix} u_1 & \dots & u_n \end{bmatrix} \begin{bmatrix} u_1 \\ \dots \\ u_n \end{bmatrix} = \sum_{i=1}^n u_i^2.$$

Alternatively, if we write $\mathbf{u} \mathbf{u}^T$, we get an **outer-product matrix**:

$$\mathbf{u} \mathbf{u}^T = \begin{bmatrix} u_1 \\ \dots \\ u_n \end{bmatrix} \begin{bmatrix} u_1 & \dots & u_n \end{bmatrix} = \begin{bmatrix} u_1^2 & \dots & u_1 u_n \\ \cdot & \dots & \cdot \\ u_n u_1 & \dots & u_n^2 \end{bmatrix}.$$

We prove the most important result of this course.

Theorem 83 (Ordinary Least Squares Estimator)

Assume that we have n observations and K explanatory variables. Suppose that we had the model:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u} \quad (33)$$

whereby $\mathbf{y}, \mathbf{u}$ are $n \times 1$ vectors, $\mathbf{X}$ is a $n \times K$ matrix, and $\boldsymbol{\beta}$ is a $K \times 1$ vector. Assume that the design matrix $\mathbf{X}$ is **full rank**. Then, the **ordinary least squares estimator** is given by:

$$\hat{\boldsymbol{\beta}}_{OLS} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}. \quad (34)$$

Proof. We have that the model is:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}$$

The goal of ordinary least squares is to minimise the residuals sum of squares by choosing a vector $\hat{\boldsymbol{\beta}}$ that does so. The residuals sum of squares can be written as:

$$\begin{aligned} \mathbf{u}^T \mathbf{u} &= [\mathbf{y} - \mathbf{X}\boldsymbol{\beta}]^T [\mathbf{y} - \mathbf{X}\boldsymbol{\beta}] \\ &= [\mathbf{y}^T - \boldsymbol{\beta}^T \mathbf{X}^T] [\mathbf{y} - \mathbf{X}\boldsymbol{\beta}] \\ &= \mathbf{y}^T \mathbf{y} - \mathbf{y}^T \mathbf{X}\boldsymbol{\beta} - \boldsymbol{\beta}^T \mathbf{X}^T \mathbf{y} + \boldsymbol{\beta}^T \mathbf{X}^T \mathbf{X}\boldsymbol{\beta} \\ &=_{(1)} \mathbf{y}^T \mathbf{y} - 2\mathbf{y}^T \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\beta}^T \mathbf{X}^T \mathbf{X}\boldsymbol{\beta} \end{aligned}$$

whereby $=_{(1)}$ uses the fact that $\boldsymbol{\beta}^T \mathbf{X}^T \mathbf{y}$ is a scalar so we have that $\boldsymbol{\beta}^T \mathbf{X}^T \mathbf{y} = (\boldsymbol{\beta}^T \mathbf{X}^T \mathbf{y})^T = \mathbf{y}^T \mathbf{X}\boldsymbol{\beta}$. To proceed further, we need some important rules from matrix calculus.

Theorem 84 (Matrix Calculus Rules)

We have the following rules of matrix calculus:

$$\frac{\partial \mathbf{a}^T \boldsymbol{\beta}}{\partial \boldsymbol{\beta}} = \mathbf{a}$$

$$\frac{\partial \boldsymbol{\beta}^T \mathbf{A} \boldsymbol{\beta}}{\partial \boldsymbol{\beta}} = 2\mathbf{A}\boldsymbol{\beta}$$

for the vectors $\mathbf{a}, \boldsymbol{\beta}$ and the symmetric matrix $\mathbf{A}$.

To minimise the residuals sum of squares, we can take the first order conditions:

$$\frac{\partial \mathbf{u}^T(\hat{\boldsymbol{\beta}}) \mathbf{u}(\hat{\boldsymbol{\beta}})}{\partial \hat{\boldsymbol{\beta}}} = \mathbf{0}$$

which then becomes:

$$\begin{aligned} \frac{\partial}{\partial \hat{\boldsymbol{\beta}}} \left[\mathbf{y}^T \mathbf{y} - 2\mathbf{y}^T \mathbf{X} \hat{\boldsymbol{\beta}} + \hat{\boldsymbol{\beta}}^T \mathbf{X}^T \mathbf{X} \hat{\boldsymbol{\beta}} \right] &= \frac{\partial \mathbf{y}^T \mathbf{y}}{\partial \hat{\boldsymbol{\beta}}} - 2 \frac{\partial \mathbf{y}^T \mathbf{X} \hat{\boldsymbol{\beta}}}{\partial \hat{\boldsymbol{\beta}}} + \frac{\partial \hat{\boldsymbol{\beta}}^T \mathbf{X}^T \mathbf{X} \hat{\boldsymbol{\beta}}}{\partial \hat{\boldsymbol{\beta}}} \\ &= \mathbf{0} - 2\mathbf{X}^T \mathbf{y} + 2\mathbf{X}^T \mathbf{X} \hat{\boldsymbol{\beta}} \end{aligned}$$

We therefore have that our first order condition is:

$$\begin{aligned} -2\mathbf{X}^T \mathbf{y} + 2\mathbf{X}^T \mathbf{X} \hat{\boldsymbol{\beta}} &= \mathbf{0} \\ \Leftrightarrow \mathbf{X}^T \mathbf{X} \hat{\boldsymbol{\beta}} &= \mathbf{X}^T \mathbf{y} \end{aligned}$$

By assumption, we assumed that $\mathbf{X}$ is full rank, which implies that $(\mathbf{X}^T \mathbf{X})$ is full rank and hence invertible, ie. $(\mathbf{X}^T \mathbf{X})^{-1}$ exists. Therefore, we have that:

$$\hat{\boldsymbol{\beta}}_{OLS} = [\mathbf{X}^T \mathbf{X}]^{-1} \mathbf{X}^T \mathbf{y}$$

□

We now look at what we have computed from the OLS estimation. We have computed that an estimate of our parameters is given by:

$$\hat{\boldsymbol{\beta}}_{OLS} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}.$$

We can use these estimated parameters to fit OLS to the data. For ease of notation, we shall rewrite $\hat{\boldsymbol{\beta}}_{OLS}$ as just $\hat{\boldsymbol{\beta}}$.

Definition 85 (Fitted Values and Residuals)

The **fitted values** are given by:

$$\hat{\mathbf{y}} = \mathbf{X} \hat{\boldsymbol{\beta}}. \quad (35)$$

The **fitted residuals** are given by:

$$\hat{\mathbf{u}} = \mathbf{y} - \hat{\mathbf{y}}. \quad (36)$$

Geometrically, we are projecting our outcome variable $\mathbf{y}$ onto the plane spanned by the explanatory variables $\mathbf{X}$. The projected point is the fitted value $\hat{\mathbf{y}} = \mathbf{X} \hat{\boldsymbol{\beta}}$. The difference between the original outcome vector $\mathbf{y}$ and this projection is the fitted residual $\hat{\mathbf{u}}$. We show this in Fig. 1.

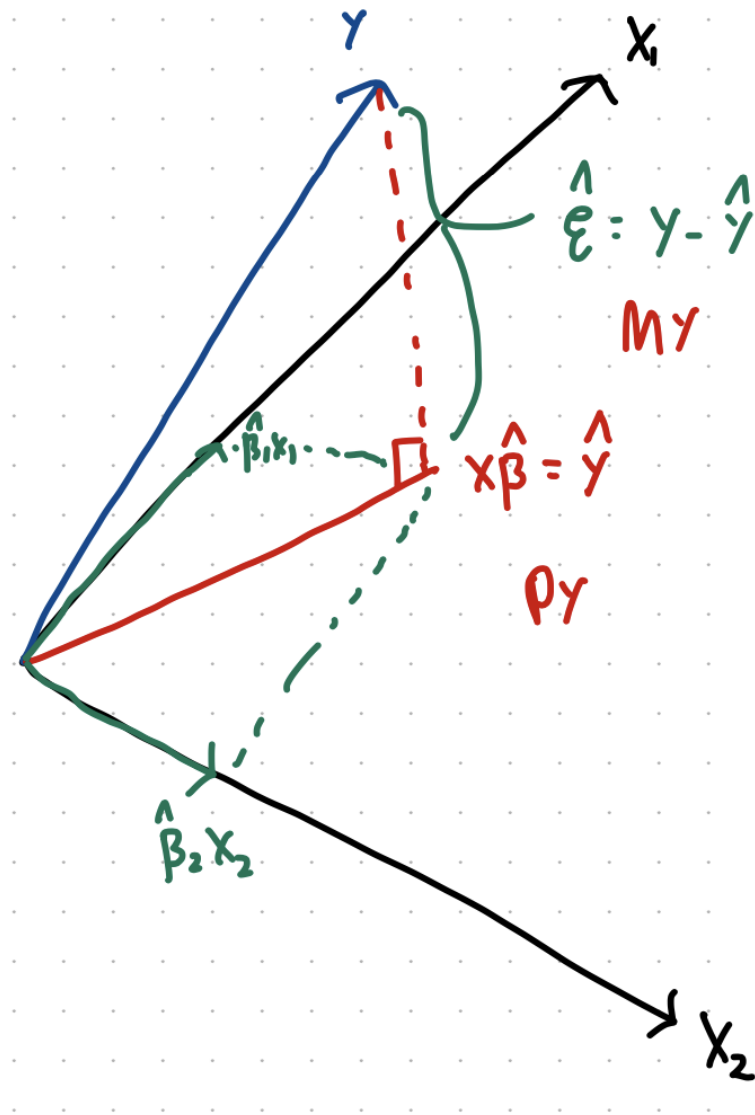


Figure 1: OLS is the projection of the observed variable onto the space spanned by the explanatory variables. This uses the example of having 2 explanatory variables.

Let us plug in our expression for the OLS estimator into our fitted value $\hat{y}$:

$$\hat{y} = X\hat{\beta} = X[(X^T X)^{-1} X^T y] = \underbrace{X(X^T X)^{-1} X^T}_P y$$

The matrix underlined by P has an important name.

Definition 86 (Projection Matrix)

The **projection matrix** P is given by:

$$P = X[X^T X]^{-1} X^T. \quad (37)$$

Remark 87. This projection matrix is a matrix such that whenever it is applied to another vector/matrix B , it will project that matrix B onto the subspace that is spanned by X .

Hence, we can write the fitted value of $\mathbf{y}$ as the projection of $\mathbf{y}$ onto the subspace spanned by $\mathbf{X}$:

$$\hat{\mathbf{y}} = \mathbf{P}\mathbf{y} \quad (38)$$

Now, if we write out fitted residuals, we get another important matrix:

$$\hat{\mathbf{u}} = \mathbf{y} - \hat{\mathbf{y}} = \mathbf{y} - \mathbf{X}[\mathbf{X}^T \mathbf{X}]^{-1} \mathbf{X}^T \mathbf{y} = \mathbf{y} - \mathbf{P}\mathbf{y} = [\mathbf{I}_n - \mathbf{P}]\mathbf{y}$$

From this, we can also define another important matrix that is thought of as the complement to the projection matrix.

Definition 88 (Annihilator Matrix)

The **annihilator matrix** is given by:

$$\mathbf{M} = \mathbf{I}_n - \mathbf{P} \quad (39)$$

Remark 89. *The way to think of the annihilator matrix is that it produces the vector that is orthogonal to the space that is spanned by $\mathbf{X}$.*

Therefore, we can write the fitted residuals to be the annihilator matrix applied to the observed variables $\mathbf{y}$.

$$\hat{\mathbf{u}} = \mathbf{M}\mathbf{y} \quad (40)$$

The following result shows where the annihilator matrix gets its name from.

Proposition 90 (Annihilator matrix is Orthogonal to Design Matrix)

The annihilator matrix is orthogonal $\mathbf{M}$ to the design matrix $\mathbf{X}$:

$$\mathbf{M} \cdot \mathbf{X} = \mathbf{0}. \quad (41)$$

Proof.

$$\begin{aligned} \mathbf{M} \cdot \mathbf{X} &= (\mathbf{I} - \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T) \mathbf{X} \\ &= \mathbf{X} - \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X} \\ &= \mathbf{X} - \mathbf{X} \\ &= \mathbf{0}. \end{aligned}$$

□

Proposition 91 (The Fitted Residuals are Orthogonal to Design Matrix)

The Fitted Residuals are Orthogonal to Design Matrix:

$$\mathbf{X}^T \hat{\mathbf{u}} = \mathbf{0} \quad (42)$$

Proof.

$$\begin{aligned}
\mathbf{X}^T \hat{\mathbf{u}} &= \mathbf{X}^T [\mathbf{M}\mathbf{y}] \\
&= \mathbf{X}^T [\mathbf{I}_n - \mathbf{P}]\mathbf{y} \\
&= \mathbf{X}^T [\mathbf{I}_n - \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T]\mathbf{y} \\
&= [\mathbf{X}^T - \mathbf{X}^T \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T]\mathbf{y} \\
&= [\mathbf{X}^T - \mathbf{I}_K \mathbf{X}^T]\mathbf{y} \\
&= [\mathbf{0}]\mathbf{y} = \mathbf{0}
\end{aligned}$$

□

We now show an important result.

Theorem 92 (Symmetry and Idempotency of $\mathbf{P}$ and $\mathbf{M}$)

The projection matrix $\mathbf{P}$ and the annihilator matrix $\mathbf{M}$ are:

1. Idempotent $\mathbf{P}\mathbf{P} = \mathbf{P}$ and $\mathbf{M}\mathbf{M} = \mathbf{M}$;
2. Symmetric $\mathbf{P}^T = \mathbf{P}$ and $\mathbf{M}^T = \mathbf{M}$.

We can also show another useful result.

Proposition 93 (Trace of the Annihilator Matrix)

The trace of the annihilator matrix $\mathbf{M}$ is given by:

$$tr(\mathbf{M}) = n - K \quad (43)$$

where n is the number of observations and K is the number of regressors.

Proof. To show this proposition, we need an important result regarding the trace operator:

Proposition 94 (Cyclic Properties of the Trace Operator)

Let $\mathbf{A}, \mathbf{B}, \mathbf{C}$ be matrices such that $\mathbf{ABC}$ is a **square matrix**. Then, the trace operator has the cyclic property:

$$tr(\mathbf{ABC}) = tr(\mathbf{CAB}) = tr(\mathbf{BCA}).$$

Now we can use this cyclic property of the trace operator to help us.

$$\begin{aligned}
tr(\mathbf{M}) &= tr[\mathbf{I}_N - \mathbf{P}] \\
&= tr[\mathbf{I}_N - \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T] \\
&= tr[\mathbf{I}_N] - tr[\mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T] \\
&=_{(1)} tr[\mathbf{I}_N] - tr[\mathbf{X}^T \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1}] \\
&=_{(2)} tr[\mathbf{I}_N] - tr[\mathbf{I}_K] \\
&= n - K
\end{aligned}$$

where $=_{(1)}$ uses the cyclic property of the trace operator and $=_{(2)}$ uses the fact that $\mathbf{X}^T \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1}$ is a $K \times K$ matrix. □

To summarise, we can think of OLS as partitioning the outcome variable $\mathbf{y}$ into two components based on the projection and the residual from the projection:

$$\mathbf{y} = \mathbf{P}\mathbf{y} + \mathbf{M}\mathbf{y}.$$

4.2 Partition Regression

Recall the formula for the OLS estimator:

$$\hat{\boldsymbol{\beta}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} \quad (44)$$

Suppose that we do not want to compute all the coefficients in $\hat{\boldsymbol{\beta}}$ but only a subset. To this end, let us partition the data matrix $\mathbf{X}$ into two sets of variables $\mathbf{X}_1$ and $\mathbf{X}_2$ whereby $\mathbf{X}_1$ is a $n \times K_1$ matrix and $\mathbf{X}_2$ is a $n \times K_2$ matrix. Without loss of generality, suppose we are interested in estimating the K_1 coefficients. Then, we have that:

$$\mathbf{y} = \begin{bmatrix} \mathbf{X}_1 & \mathbf{X}_2 \end{bmatrix} \begin{bmatrix} \boldsymbol{\beta}_1 \\ \boldsymbol{\beta}_2 \end{bmatrix} + \boldsymbol{\epsilon} = \mathbf{X}_1 \boldsymbol{\beta}_1 + \mathbf{X}_2 \boldsymbol{\beta}_2 + \boldsymbol{\epsilon} \quad (45)$$

Here, we are interested in estimating only $\boldsymbol{\beta}_1$. We know that the first order condition for OLS is given by:

$$\mathbf{X}^T \mathbf{y} = \mathbf{X}^T \mathbf{X} \boldsymbol{\beta}$$

Therefore, substituting our partitioned matrices into the first order condition gives us:

$$\begin{bmatrix} \mathbf{X}_1^T \mathbf{y} \\ \mathbf{X}_2^T \mathbf{y} \end{bmatrix} = \begin{bmatrix} \mathbf{X}_1^T \mathbf{X}_1 & \mathbf{X}_1^T \mathbf{X}_2 \\ \mathbf{X}_2^T \mathbf{X}_1 & \mathbf{X}_2^T \mathbf{X}_2 \end{bmatrix} \begin{bmatrix} \boldsymbol{\beta}_1 \\ \boldsymbol{\beta}_2 \end{bmatrix}$$

We can then solve for our OLS parameters:

$$\begin{bmatrix} \hat{\boldsymbol{\beta}}_1 \\ \hat{\boldsymbol{\beta}}_2 \end{bmatrix} = \begin{bmatrix} \mathbf{X}_1^T \mathbf{X}_1 & \mathbf{X}_1^T \mathbf{X}_2 \\ \mathbf{X}_2^T \mathbf{X}_1 & \mathbf{X}_2^T \mathbf{X}_2 \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{X}_1^T \mathbf{y} \\ \mathbf{X}_2^T \mathbf{y} \end{bmatrix}. \quad (46)$$

We are interested in figuring out what $\hat{\boldsymbol{\beta}}_1$ is. The first block of normal equations gives:

$$\mathbf{X}_1^T \mathbf{X}_1 \hat{\boldsymbol{\beta}}_1 + \mathbf{X}_1^T \mathbf{X}_2 \hat{\boldsymbol{\beta}}_2 = \mathbf{X}_1^T \mathbf{y},$$

so

$$\hat{\boldsymbol{\beta}}_1 = (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T (\mathbf{y} - \mathbf{X}_2 \hat{\boldsymbol{\beta}}_2).$$

Now, suppose that $\mathbf{X}_1$ and $\mathbf{X}_2$ are **orthogonal** to each other so that:

$$\mathbf{X}_1^T \mathbf{X}_2 = \mathbf{0}.$$

Then, we have that our estimator for $\boldsymbol{\beta}_1$ becomes:

$$\hat{\boldsymbol{\beta}}_1 = (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{y}$$

This has a very nice interpretation that if we want to get the estimators on the K_1 coefficients, we only need to regress $\mathbf{y}$ on the K_1 explanatory variables in $\mathbf{X}_1$ and we are done!

However, what if $\mathbf{X}_1$ and $\mathbf{X}_2$ are not orthogonal? Then we need to rely on the Frisch-Waugh-Lovell theorem in the next section.

4.3 Frisch-Waugh-Lovell Theorem

So far, we have arrived at the expression for the OLS parameters:

$$\begin{bmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \end{bmatrix} = \begin{bmatrix} \mathbf{X}_1^T \mathbf{X}_1 & \mathbf{X}_1^T \mathbf{X}_2 \\ \mathbf{X}_2^T \mathbf{X}_1 & \mathbf{X}_2^T \mathbf{X}_2 \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{X}_1^T \mathbf{y} \\ \mathbf{X}_2^T \mathbf{y} \end{bmatrix}. \quad (47)$$

Suppose that we want to derive an expression for $\hat{\beta}_2$. We can do this in a two step procedure.

Step 1: Regress $\mathbf{X}_2$ on $\mathbf{X}_1$

The first step is to regress $\mathbf{X}_2$ onto $\mathbf{X}_1$ and retrieve the residuals $\mathbf{X}_{2,1}$:

$$\mathbf{X}_2 = \mathbf{X}_1 \mathbf{b} + \mathbf{u} \quad (48)$$

where $\mathbf{b}$ is a $K_1 \times K_2$ matrix of parameters. This gives us the OLS estimator:

$$\hat{\mathbf{b}} = (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_2$$

where we project $\mathbf{X}_2$ onto the subspace spanned by $\mathbf{X}_1$. We can then retrieve the residuals $\mathbf{X}_{2,1}$ as the true value $\mathbf{X}_2$ less the OLS projection:

$$\mathbf{X}_{2,1} = \mathbf{X}_2 - \mathbf{X}_1 \hat{\mathbf{b}}$$

What we have done by regression $\mathbf{X}_2$ on $\mathbf{X}_1$ and retrieving the residuals is that we retrieved the variation in $\mathbf{X}_2$ that is **not explained** by the variation in $\mathbf{X}_1$. This can be thought of as purging the effects of $\mathbf{X}_1$ on $\mathbf{X}_2$.

Step 2: Regress $\mathbf{y}$ on $\mathbf{X}_{2,1}$

We now regress $\mathbf{y}$ on the residuals $\mathbf{X}_{2,1}$ and this gets us our desired estimator $\hat{\beta}_2$. We now show why this would work. First, we can write out:

$$\begin{aligned} \mathbf{y} &= \mathbf{X}_1 \boldsymbol{\beta}_1 + \mathbf{X}_2 \boldsymbol{\beta}_2 + \boldsymbol{\epsilon} \\ &=_{(1)} \mathbf{X}_1 \boldsymbol{\beta}_1 + \mathbf{X}_2 \boldsymbol{\beta}_2 + \boldsymbol{\epsilon} + \mathbf{X}_1 \hat{\mathbf{b}} \boldsymbol{\beta}_2 - \mathbf{X}_1 \hat{\mathbf{b}} \boldsymbol{\beta}_2 \\ &= \mathbf{X}_1 (\boldsymbol{\beta}_1 + \hat{\mathbf{b}} \boldsymbol{\beta}_2) + (\mathbf{X}_2 - \mathbf{X}_1 \hat{\mathbf{b}}) \boldsymbol{\beta}_2 + \boldsymbol{\epsilon} \\ &=_{(2)} \mathbf{X}_1 \boldsymbol{\delta}_1 + \mathbf{X}_{2,1} \boldsymbol{\beta}_2 + \boldsymbol{\epsilon} \end{aligned}$$

where $=_{(1)}$ involved $\pm \mathbf{X}_1 \hat{\mathbf{b}} \boldsymbol{\beta}_2$ and $=_{(2)}$ involved defining $\boldsymbol{\delta}_1 \equiv \boldsymbol{\beta}_1 + \hat{\mathbf{b}} \boldsymbol{\beta}_2$ and recognising that $\mathbf{X}_{2,1} = \mathbf{X}_2 - \mathbf{X}_1 \hat{\mathbf{b}}$ are the residuals from regression $\mathbf{X}_2$ on $\mathbf{X}_1$.

Now, we again use our first order conditions to see that:

$$\begin{bmatrix} \hat{\boldsymbol{\delta}}_1 \\ \hat{\boldsymbol{\beta}}_2 \end{bmatrix} = \begin{bmatrix} \mathbf{X}_1^T \mathbf{X}_1 & \mathbf{X}_1^T \mathbf{X}_{2,1} \\ \mathbf{X}_{2,1}^T \mathbf{X}_1 & \mathbf{X}_{2,1}^T \mathbf{X}_{2,1} \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{X}_1^T \mathbf{y} \\ \mathbf{X}_{2,1}^T \mathbf{y} \end{bmatrix}$$

From this, we can state the expression for our desired estimator $\hat{\beta}_2$.

Theorem 95 (Coefficients of Partition Regression)

The coefficients of the partition regression is given by:

$$\begin{bmatrix} \hat{\delta}_1 \\ \hat{\beta}_2 \end{bmatrix} = \begin{bmatrix} (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{y} \\ (\mathbf{X}_{2,1}^T \mathbf{X}_{2,1})^{-1} \mathbf{X}_{2,1}^T \mathbf{y} \end{bmatrix} \quad (49)$$

Proof. Recall that we had:

$$\begin{bmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \end{bmatrix} = \begin{bmatrix} \mathbf{X}_1^T \mathbf{X}_1 & \mathbf{X}_1^T \mathbf{X}_2 \\ \mathbf{X}_2^T \mathbf{X}_1 & \mathbf{X}_2^T \mathbf{X}_2 \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{X}_1^T \mathbf{Y} \\ \mathbf{X}_2^T \mathbf{Y} \end{bmatrix}$$

We can replace:

$$\begin{bmatrix} \hat{\delta}_1 \\ \hat{\beta}_2 \end{bmatrix} = \begin{bmatrix} \mathbf{X}_1^T \mathbf{X}_1 & \mathbf{X}_1^T \mathbf{X}_{2,1} \\ \mathbf{X}_{2,1}^T \mathbf{X}_1 & \mathbf{X}_{2,1}^T \mathbf{X}_{2,1} \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{X}_1^T \mathbf{Y} \\ \mathbf{X}_{2,1}^T \mathbf{Y} \end{bmatrix}$$

Now, we want to show that the off-diagonal elements are zero.

Proposition 96 (Nuisance Vector and Residuals from First Stage Are Orthogonal)

$\mathbf{X}_1$ and $\mathbf{X}_{2,1}$ are orthogonal:

$$\mathbf{X}_1^T \cdot \mathbf{X}_{2,1} = \mathbf{0} \quad (50)$$

Proof.

$$\begin{aligned} \mathbf{X}_1^T \cdot \mathbf{X}_{2,1} &= \mathbf{X}_1^T \cdot (\mathbf{X}_2 - \mathbf{X}_1 \hat{\mathbf{b}}) \\ &= \mathbf{X}_1^T \cdot (\mathbf{X}_2 - \mathbf{X}_1 (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_2) \\ &= \mathbf{X}_1^T \mathbf{X}_2 - \mathbf{X}_1^T \mathbf{X}_1 (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_2 \\ &= \mathbf{X}_1^T \mathbf{X}_2 - \mathbf{X}_1^T \mathbf{X}_2 \\ &= \mathbf{0} \end{aligned}$$

□

Therefore, we have just shown that $\mathbf{X}_1^T \cdot \mathbf{X}_{2,1} = \mathbf{0}$ and so we have:

$$\begin{bmatrix} \hat{\delta}_1 \\ \hat{\beta}_2 \end{bmatrix} = \begin{bmatrix} \mathbf{X}_1^T \mathbf{X}_1 & 0 \\ 0 & \mathbf{X}_{2,1}^T \mathbf{X}_{2,1} \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{X}_1^T \mathbf{Y} \\ \mathbf{X}_{2,1}^T \mathbf{Y} \end{bmatrix}$$

which gives us our desired result. □

Therefore, we can simply read off the second line in the theorem we have just proved to see our expression for our desired estimator as being:

$$\hat{\beta}_2 = (\mathbf{X}_{2,1}^T \mathbf{X}_{2,1})^{-1} \mathbf{X}_{2,1}^T \mathbf{Y}$$

This two-step procedure of regressing the variable of interest on the nuisance variables and then regressing the outcome on the residuals is known as the **Frisch-Waugh-Lovell theorem**.

4.4 Residual Variance Decomposition

Now, we want to derive another useful relationship.

Theorem 97 (Partial Sum of Squares Decomposition)

Partition $\mathbf{X} = [\mathbf{X}_1 \ \mathbf{X}_2]$. Then:

Restricted-model RSS = Incremental ESS from $\mathbf{X}_2$ + Unrestricted-model RSS.

If $\mathbf{X}_1$ contains only an intercept, the restricted-model RSS is the usual centred total sum of squares.

Proof. First, start off with the fact that the fitted residuals are:

$$\hat{\mathbf{u}} = \mathbf{M}\mathbf{y} = (\mathbf{I} - \mathbf{P})\mathbf{y} = (\mathbf{I} - \mathbf{X}(\mathbf{X}^T\mathbf{X})^{-1}\mathbf{X}^T)\mathbf{y}$$

Then, the residual sum of squares is:

$$\hat{\mathbf{u}}^T \hat{\mathbf{u}} = ((\mathbf{I} - \mathbf{P})\mathbf{y})^T ((\mathbf{I} - \mathbf{P})\mathbf{y}) = \mathbf{y}^T (\mathbf{I} - \mathbf{P})^T (\mathbf{I} - \mathbf{P}) \mathbf{y} = \mathbf{y}^T (\mathbf{I} - \mathbf{P}) \mathbf{y}$$

where we used the symmetry and idempotent of $\mathbf{M}$.

Now, partition the equation $\mathbf{y} = \mathbf{X}_1\boldsymbol{\delta}_1 + \mathbf{X}_{2,1}\boldsymbol{\beta}_2 + \mathbf{u}$ whereby $\mathbf{X}_1^T \mathbf{X}_{2,1} = \mathbf{0}$. Re-arrange this to get:

$$\hat{\mathbf{u}} = \mathbf{y} - \mathbf{X}_1\hat{\boldsymbol{\delta}}_1 - \mathbf{X}_{2,1}\hat{\boldsymbol{\beta}}_2$$

Now, plug in our expression for $\hat{\boldsymbol{\delta}}_1$ and $\hat{\boldsymbol{\beta}}_2$ we derived before to get:

$$\hat{\mathbf{u}} = \mathbf{y} - \underbrace{\mathbf{X}_1(\mathbf{X}_1^T\mathbf{X}_1)^{-1}\mathbf{X}_1^T}_{\equiv \mathbf{P}_1} \mathbf{y} - \underbrace{\mathbf{X}_{2,1}(\mathbf{X}_{2,1}^T\mathbf{X}_{2,1})^{-1}\mathbf{X}_{2,1}^T}_{\equiv \mathbf{P}_{2,1}} \mathbf{y}$$

and this becomes

$$\hat{\mathbf{u}} = \mathbf{y} - \mathbf{P}_1\mathbf{y} - \mathbf{P}_{2,1}\mathbf{y} = (\mathbf{I} - \mathbf{P}_1 - \mathbf{P}_{2,1})\mathbf{y}$$

Then, this means that our residual sum of squares becomes:

$$\begin{aligned} \hat{\mathbf{u}}^T \hat{\mathbf{u}} &= ((\mathbf{I} - \mathbf{P}_1 - \mathbf{P}_{2,1})\mathbf{y})^T ((\mathbf{I} - \mathbf{P}_1 - \mathbf{P}_{2,1})\mathbf{y}) \\ &= \mathbf{y}^T (\mathbf{I} - \mathbf{P}_1 - \mathbf{P}_{2,1})^T (\mathbf{I} - \mathbf{P}_1 - \mathbf{P}_{2,1}) \mathbf{y} \\ &= \mathbf{y}^T (\mathbf{I} - \mathbf{P}_1 - \mathbf{P}_{2,1}) \mathbf{y} \\ &= \mathbf{y}^T (\mathbf{I} - \mathbf{P}_1) \mathbf{y} - \mathbf{y}^T \mathbf{P}_{2,1} \mathbf{y} \end{aligned}$$

So we then have that:

$$\underbrace{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}_{\text{unrestricted RSS}} = \underbrace{\mathbf{y}^T (\mathbf{I} - \mathbf{P}_1) \mathbf{y}}_{\text{restricted RSS}} - \underbrace{\mathbf{y}^T \mathbf{P}_{2,1} \mathbf{y}}_{\text{incremental ESS}}$$

Therefore, restricted RSS equals incremental ESS plus unrestricted RSS:

$$\mathbf{y}^T (\mathbf{I} - \mathbf{P}_1) \mathbf{y} = \mathbf{y}^T \mathbf{P}_{2,1} \mathbf{y} + \mathbf{y}^T (\mathbf{I} - \mathbf{P}_1 - \mathbf{P}_{2,1}) \mathbf{y}$$

□

5 Asymptotic Theory

5.1 Convergence in Distribution and Probability

We now state two important convergence concepts we will use frequently.

Definition 98 (Convergence in Distribution)

A sequence of random variables $X_1, X_2, \dots$ **converges in distribution** to a random variable X if

$$\lim_{n \rightarrow \infty} F_{X_n}(x) = F_X(x) \quad (51)$$

at all points x where $F_X(x)$ is continuous, where F_{X_n} is the distribution function of X_n and F_X is the distribution function of X .

Remark 99. The distribution to which the sequence X_n converges to is referred to as the **limit distribution**.

Definition 100 (Convergence in probability)

A sequence of random variables $X_1, X_2, \dots$ **converges in probability** to a random variable X if for every $\epsilon > 0$

$$\lim_{n \rightarrow \infty} P[|X_n - X| \geq \epsilon] = 0$$

Proposition 101 (Convergence in probability implies convergence in distribution)

Convergence in probability implies convergence in distribution. That is, if

$$X_n \xrightarrow{P} X$$

then

$$X_n \xrightarrow{D} X$$

Example 102

Let $X_1, X_2, \dots$ be a sequence of random variables and suppose that X has a standard normal distribution. If

$$X_n \xrightarrow{P} X$$

By our proposition, this means that

$$X_n \xrightarrow{D} X$$

Therefore, the limiting random variable has the standard normal probability density function

$$f_X(x) = \frac{1}{\sqrt{2\pi}} e^{-x^2/2}.$$

Proposition 103

The sequence of random variables $X_1, X_2, \dots$ converges in probability to a constant μ if and only if the sequence converges in distribution to the constant μ .

Definition 104 (Consistent Estimator)

An estimator $\hat{\theta}_n$ is a consistent estimator of the parameter θ if:

$$\hat{\theta}_n \xrightarrow{P} \theta$$

Consistency is an asymptotic property. This is in contrast to unbiasedness, which is a finite-sample property.

Consistency does not imply unbiasedness and vice versa. Consistency means that the distribution of $\hat{\theta}_n$ concentrates arbitrarily closely around θ as $n \rightarrow \infty$; convergence of the variance to zero is a useful sufficient condition when the relevant moments exist, but it is not part of the definition.

Example 105

Let $\hat{\theta}$ be an unbiased and consistent estimator of θ . Define the estimator:

$$\tilde{\theta} = \hat{\theta} + \frac{1}{n}$$

Clearly, this new estimator is biased:

$$\mathbb{E}[\tilde{\theta}] = \mathbb{E}[\hat{\theta} + \frac{1}{n}] = \theta + \frac{1}{n} \neq \theta.$$

However, it is consistent:

$$p \lim_{n \rightarrow \infty} \tilde{\theta} = p \lim_{n \rightarrow \infty} \hat{\theta} + p \lim_{n \rightarrow \infty} \frac{1}{n} = \theta + 0 = \theta.$$

5.1.1 Weak Law of Large Numbers

There are two versions of the weak law of large numbers we will need.

Theorem 106 (Weak Law of Large Numbers (1))

Let $X_1, X_2, \dots$ be **independent** random variables with the conditions:

1. $\mathbb{E}[X_i] = \mu < \infty$
2. $\text{Var}[X_i] = \sigma^2 < \infty$

Then:

$$\frac{1}{n} \sum_{i=1}^n X_i = \bar{X}_n \xrightarrow{P} \mathbb{E}[X_i] = \mu.$$

Alternatively, we can write that:

$$p \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n X_i = \mathbb{E}[X_i] = \mu.$$

Theorem 107 (Weak Law of Large Numbers (2))

Let $X_1, X_2, \dots$ be **independent and identically distributed** random variables with the condition:

$$1. \mathbb{E}[X_i] = \mu < \infty$$

Then:

$$\frac{1}{n} \sum_{i=1}^n X_i = \bar{X}_n \xrightarrow{P} \mathbb{E}[X_i] = \mu.$$

Alternatively, we can write that:

$$P \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n X_i = \mathbb{E}[X_i] = \mu.$$

Here, notice that version (1) of the weak law of large numbers (WLLN) required only independent random variables but we then required that the variance is finite. Meanwhile, in version (2), we required the random variables to also be identically distributed but then we only required that the first moment is finite. Whichever one you use depends on what is given to you.

Example 108

(WLLN Applied to Bernoulli Random Variables). Let $X_1, \dots, X_n$ be i.i.d **Bernoulli random variables** with mean p . Then, we can apply the WLLN (2) to conclude that:

$$\frac{1}{n} \sum_{i=1}^n X_i \xrightarrow{P} \mathbb{E}[X_i] = p.$$

Separately, linearity of expectations gives $\mathbb{E}[\sum_{i=1}^n X_i] = np$; indeed, the sum has a **Binomial distribution**.

We show a very useful application of the weak law of large numbers to show a useful property of an important estimator.

Proposition 109 (Population Variance Estimator is Consistent)

Let $X_1, X_2, \dots$ be i.i.d random variables with mean μ and variance σ^2 . Then, the estimator of the variance:

$$\hat{\sigma}_n^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \mu)^2 \tag{52}$$

is a **consistent estimator** of σ^2 .

Proof. First, we take the expectation of the terms inside the brackets:

$$\mathbb{E}[X_i - \mu] = \mathbb{E}[X_i] - \mu = \mu - \mu = 0.$$

Now, we know that for any random variable Z , the variance of Z is given by:

$$\text{Var}[Z] = \mathbb{E}[Z^2] - \mathbb{E}[Z]^2$$

but if we have that the first moment of Z is $\mathbb{E}[Z] = 0$, then

$$\text{Var}[Z] = \mathbb{E}[Z^2].$$

We have shown that this is indeed the case for $X_i - \mu$ so we therefore have that:

$$\text{Var}[(X_i - \mu)] = \mathbb{E}[(X_i - \mu)^2] - \mathbb{E}[(X_i - \mu)]^2 = \mathbb{E}[(X_i - \mu)^2] - 0 = \mathbb{E}[(X_i - \mu)^2]$$

We now apply the weak law of large numbers (2) to our estimator:

$$\begin{aligned}\hat{\sigma}_n^2 &= \frac{1}{n} \sum_{i=1}^n (X_i - \mu)^2 \xrightarrow{P} \mathbb{E}[(X_i - \mu)^2] \\ &=_{(1)} \text{Var}[X_i - \mu] \\ &=_{(2)} \text{Var}[X_i] \\ &= \sigma^2\end{aligned}$$

whereby $=_{(1)}$ comes from what we have shown earlier and $=_{(2)}$ comes from the fact that μ is a constant that is added and hence can be dropped from the variance. Therefore, we have shown that:

$$\hat{\sigma}_n^2 \xrightarrow{P} \sigma^2$$

or equivalently:

$$p \lim_{n \rightarrow \infty} \hat{\sigma}_n^2 = \sigma^2$$

Therefore, $\hat{\sigma}_n^2$ is a **consistent** estimator. □

Remark 110. In practice, we do not know μ and need to estimate it with $\hat{\mu}$. Replacing μ with $\hat{\mu}$ therefore requires an additional argument, which we give below.

5.1.2 Continuous Mapping Theorem

Theorem 111 (Continuous Mapping Theorem)

Suppose that $X_1, X_2, \dots$ converges in probability to a random variable X

$$X_n \xrightarrow{P} X$$

Let $h(\cdot)$ be a continuous function. Then, the **continuous mapping theorem** states that we have that:

$$h(X_n) \xrightarrow{P} h(X)$$

This also holds for convergence in distribution $\xrightarrow{D}$.

Example 112

(Consistency of standard deviation estimator). The sample variance S_n^2 is a consistent estimator for the variance σ^2 . Define the continuous function $h(x) = \sqrt{x}$. Then,

$$S_n = \sqrt{S_n^2} = h(S_n^2)$$

is a consistent estimator of the standard deviation σ .

The continuous mapping theorem gives us some useful results.

Theorem 113 (Implications of Continuous Mapping Theorem)

Suppose that:

$$X_n \xrightarrow{P} X \quad \text{and} \quad Y_n \xrightarrow{P} Y.$$

Then:

1. $X_n + Y_n \xrightarrow{P} X + Y;$
2. $X_n \cdot Y_n \xrightarrow{P} X \cdot Y.$

Theorem 114 (Slutsky's Theorem)

Let $\{X_n\}_{n=1}^{\infty}$ and $\{Y_n\}_{n=1}^{\infty}$ be sequences of random variables. Let α be a constant. Suppose that $X_n \xrightarrow{D} X$ and $Y_n \xrightarrow{P} \alpha$. Then

1. $Y_n X_n \xrightarrow{D} \alpha X$
2. $X_n + Y_n \xrightarrow{D} X + \alpha$

Example 115

Let $X_1, X_2, \dots$ be standard normal random variables with a limit $X \sim \mathcal{N}(0, 1)$. Let $g(x) = x^2$. Then

$$g(X_i) = X_i^2 \xrightarrow{D} \chi^2(1).$$

Example 116

(Normal approximation with estimated variance). Suppose that we have

$$\frac{\sqrt{n}(\bar{X} - \mu)}{\sigma} \xrightarrow{D} \mathcal{N}(0, 1)$$

where σ is unknown. We can plug in our sample variance estimator S_n^2 . We are interested in the distribution of:

$$\frac{\sqrt{n}(\bar{X} - \mu)}{S_n}$$

We know that the sample variance estimator is a consistent estimator, so we have that:

$$\frac{\sigma}{S_n} \xrightarrow{P} \frac{\sigma}{\sigma} = 1.$$

Then, by multiplying and dividing by σ , and invoking **Slutsky's theorem**, we have that

$$\frac{\sqrt{n}(\bar{X} - \mu)}{S_n} = \frac{\sigma}{S_n} \frac{\sqrt{n}(\bar{X} - \mu)}{\sigma} \xrightarrow{D} \mathcal{N}(0, 1).$$

Proposition 117 (Sample Variance Estimator is Consistent)

Let $X_1, X_2, \dots$ be i.i.d random variables with mean μ and variance σ^2 . Then, the estimator of the variance:

$$\hat{\sigma}_n^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 \tag{53}$$

is a **consistent estimator** of σ^2 .

Proof. First, we can rewrite the sample variance estimator as:

$$\hat{\sigma}_n^2 = \underbrace{\frac{n}{n-1}}_{(A)} \left(\underbrace{\frac{1}{n} \sum_{i=1}^n X_i^2}_{(B)} - \underbrace{\left(\frac{1}{n} \sum_{i=1}^n X_i \right)^2}_{(C)} \right)$$

We now analyse each of these terms.

For (A), we have that clearly:

$$\frac{n}{n-1} = \frac{1}{1-1/n} \rightarrow 1.$$

Now for (B), recall that $\text{Var}[X] = \mathbb{E}[X^2] - \mathbb{E}[X]^2$. We have that:

$$\mathbb{E}[X^2] = \sigma^2 + \mu^2$$

if $\mathbb{E}[X] = \mu$ and $\text{Var}[X] = \sigma^2$.

Now, we know that $X_1, \dots, X_n$ are i.i.d and $g(x) = x^2$ is a continuous function. This implies that $g(X_i) = X_i^2$ is also i.i.d. Hence, we have a sequence of i.i.d random variables with finite first moment. We can apply the **weak law of large numbers** to (B) to get:

$$\frac{1}{n} \sum_{i=1}^n X_i^2 \xrightarrow{P} \mathbb{E}[X_i^2] = \sigma^2 + \mu^2.$$

We now look at (C). We have that for the term inside the square, we can again apply the **weak law of large numbers**:

$$\frac{1}{n} \sum_{i=1}^n X_i \xrightarrow{P} \mu.$$

Then, we can apply the **continuous mapping theorem**

$$\left(\frac{1}{n} \sum_{i=1}^n X_i \right)^2 \xrightarrow{P} (\mu)^2 = \mu^2.$$

We now bring everything together. We first notice that $f(a, b) = a - b$ is a continuous function, so we can apply the continuous mapping theorem to conclude:

$$\hat{\sigma}_n^2 = \underbrace{\frac{n}{n-1}}_{(A)} \left(\underbrace{\frac{1}{n} \sum_{i=1}^n X_i^2}_{(B)} - \underbrace{\left(\frac{1}{n} \sum_{i=1}^n X_i \right)^2}_{(C)} \right) \xrightarrow{P} 1 \left[(\sigma^2 + \mu^2) - \mu^2 \right] = \sigma^2$$

Therefore, we have shown that:

$$p \lim_{n \rightarrow \infty} \hat{\sigma}_n^2 = \sigma^2$$

is a consistent estimator. □

5.1.3 Central Limit Theorem

The central limit theorem states that the standardised mean, rescaled by $\sqrt{n}$, converges in distribution to a standard normal. That is, if $X_1, \dots, X_n$ are i.i.d random variables with finite mean $\mathbb{E}[X_i] = \mu$ and variance $\text{Var}[X_i] = \sigma^2 > 0$, then:

$$\left[\frac{\bar{X}_n - \mathbb{E}[\bar{X}_n]}{\sqrt{\text{Var}[\bar{X}_n]}} \right] \xrightarrow{D} \mathcal{N}(0, 1) \tag{54}$$

whereby $\bar{X}_n$ is the sample average. However, we know that:

$$\mathbb{E}[\bar{X}_n] = \mu$$

and

$$\text{Var}[\bar{X}_n] = \frac{\sigma^2}{n}$$

Therefore, we can rewrite the central limit theorem as:

$$\left[\frac{\bar{X}_n - \mu}{\sqrt{\sigma^2/n}} \right] \xrightarrow{D} \mathcal{N}(0, 1)$$

This gives us the more familiar form of the central limit theorem.

Theorem 118 (Lindeberg–Lévy Central Limit Theorem)

Let $X_1, \dots, X_n$ be i.i.d random variables with finite mean $\mathbb{E}[X_i] = \mu$ and variance $\text{Var}[X_i] = \sigma^2 > 0$. Then, the standardised sample average has the following limiting distribution:

$$\sqrt{n} \left[\frac{\bar{X}_n - \mu}{\sigma} \right] \xrightarrow{D} \mathcal{N}(0, 1). \quad (55)$$

We also state the central limit theorem for random vectors.

Theorem 119 (Central Limit Theorem for Random Vectors)

Suppose that $\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_n$ are independent and identically distributed random vectors with $\mathbb{E}[\mathbf{z}_i] = \boldsymbol{\mu}$ and finite covariance matrix $\text{Var}[\mathbf{z}_i] = \boldsymbol{\Sigma}$. Denote $\bar{\mathbf{z}}_n = \frac{1}{n} \sum_{i=1}^n \mathbf{z}_i$ as the vector of sample means for a sample of size n . Then, the central limit theorem states that:

$$\sqrt{n} [\bar{\mathbf{z}}_n - \boldsymbol{\mu}] \xrightarrow{D} \mathcal{N}(\mathbf{0}, \boldsymbol{\Sigma}). \quad (56)$$

Example 120

(Asymptotic Normality of t-statistic) Define the t-test statistic:

$$t_n = \frac{\sqrt{n}(\bar{X}_n - \mu)}{\sqrt{\hat{\sigma}^2}}$$

In the numerator, we have the **central limit theorem**:

$$\sqrt{n}(\bar{X}_n - \mu) \xrightarrow{D} \mathcal{N}(0, \sigma^2).$$

In the denominator, we know that the sample variance estimator is consistent:

$$\hat{\sigma}^2 \xrightarrow{P} \sigma^2.$$

By the **continuous mapping theorem**

$$\frac{1}{\sqrt{\hat{\sigma}^2}} \xrightarrow{P} \frac{1}{\sqrt{\sigma^2}}$$

Now, we can apply **Slutsky's theorem** to the t-test statistic to get:

$$t_n = \frac{\sqrt{n}(\bar{X}_n - \mu)}{\sqrt{\hat{\sigma}^2}} \xrightarrow{D} \frac{1}{\sigma} \mathcal{N}(0, \sigma^2) = \mathcal{N}(0, 1)$$

which only holds if $\sigma^2 > 0$ due to the continuous mapping theorem.

5.1.4 Delta Method

Sometimes we are interested in the distribution of a **function** of a random variable. The delta method tells us the limiting distribution of $g(X_n)$. For intuition, we take a first-order Taylor approximation of $g(X_n)$ and show that the remainder is asymptotically negligible.

Theorem 121 (Delta Method)

Let X_n be a sequence of random variables such that

$$\sqrt{n}(X_n - \mu) \xrightarrow{D} \mathcal{N}(0, \sigma^2) \tag{57}$$

Suppose that we have a continuously differentiable function $g : \mathbb{R} \rightarrow \mathbb{R}$. Then, the **delta method** states that:

$$\sqrt{n}(g(X_n) - g(\mu)) \xrightarrow{D} \mathcal{N}(0, \sigma^2 [g'(\mu)]^2) \tag{58}$$

provided that $g'(\mu)$ exists.

Proof. (Sketch). Differentiability at μ gives the first-order expansion

$$g(X_n) - g(\mu) = g'(\mu)(X_n - \mu) + o_p(|X_n - \mu|).$$

Multiplying by $\sqrt{n}$ and using the assumed limiting distribution yields

$$\sqrt{n}(g(X_n) - g(\mu)) = g'(\mu)\sqrt{n}(X_n - \mu) + o_p(1).$$

Slutsky's theorem therefore gives

$$\sqrt{n}(g(X_n) - g(\mu)) \xrightarrow{D} \mathcal{N}(0, \sigma^2 [g'(\mu)]^2).$$

□

An alternative way to think about the delta method is to think that $g(X_n)$ is asymptotically normal:

$$g(X_n) \overset{\Delta}{\sim} \mathcal{N}\left(g(\mu), \frac{\sigma^2 [g'(\mu)]^2}{n}\right)$$

However, it is important to note that this is an approximation as it uses a Taylor series approximation!!!

Example 122

(Squared Sample Average). Suppose $X_1, \dots, X_n$ are i.i.d with $\mathbb{E}[X_i] = \mu$ and $\text{Var}[X_i] = \sigma^2$. We are interested in the limiting distribution of $\bar{X}_n^2$. By the **central limit theorem**

$$\sqrt{n}(\bar{X}_n - \mu) \xrightarrow{D} \mathcal{N}(0, \sigma^2).$$

Define $g(y) = y^2$ so that:

$$g(\bar{X}_n) = \bar{X}_n^2$$

and

$$g(\mu) = \mu^2 \Rightarrow g'(\mu) = 2\mu$$

This means that

$$(g'(\mu))^2 = 4\mu^2.$$

We can therefore apply the delta method formula to conclude that:

$$\sqrt{n}(\bar{X}_n^2 - \mu^2) \xrightarrow{D} \mathcal{N}(0, 4\mu^2\sigma^2).$$

However, at $\mu = 0$, we must treat this differently as we get a degenerate distribution $\mathcal{N}(0, 0)$.

First, we notice for $\mu = 0$,

$$\sqrt{n}(\bar{X}_n - 0) = \sqrt{n}\bar{X}_n \xrightarrow{D} \mathcal{N}(0, \sigma^2) = \sigma\mathcal{N}(0, 1)$$

by the central limit theorem. Therefore, we have that:

$$(\sqrt{n}\bar{X}_n)^2 = n\bar{X}_n^2 \xrightarrow{D} (\sigma\mathcal{N}(0, 1))^2 = \sigma^2\chi^2(1).$$

Example 123

(Average of Poisson Random Variables). Let $X_1, \dots, X_n$ be i.i.d $Pois(\lambda)$. Define

$$\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$$

We know $\mathbb{E}[X_i] = Var[X_i] = \lambda$.

By **central limit theorem**

$$\sqrt{n} \left[\frac{\bar{X}_n - \lambda}{\sqrt{\lambda}} \right] \xrightarrow{D} \mathcal{N}(0, 1)$$

or

$$\sqrt{n}(\bar{X}_n - \lambda) \xrightarrow{D} \mathcal{N}(0, \lambda).$$

Apply $g(x) = \sqrt{x}$. We have that:

$$g(\bar{X}_n) = \sqrt{\bar{X}_n} \quad g(\lambda) = \sqrt{\lambda}$$

Then we have

$$g'(\lambda) = \frac{\lambda^{-1/2}}{2}$$

and

$$(g'(\lambda))^2 = \frac{\lambda^{-1}}{4}$$

Then, by delta theorem:

$$\sqrt{n}(\sqrt{\bar{X}_n} - \sqrt{\lambda}) \xrightarrow{D} \frac{\lambda^{-1/2}}{2} \mathcal{N}(0, \lambda) = \mathcal{N}(0, \frac{1}{4\lambda}) = \mathcal{N}(0, \frac{1}{4}).$$

This is known as a variance stabilising transformation as the variance no longer depends on λ .

Remark 124. Note that we are saying that for every possible parameter value θ , our estimator $\hat{\theta}$ needs to converge in probability to the true parameter θ .

5.2 Stochastic Order

We now are interested in formalising the ideas we have talked about so far regarding convergence of random variables. What we will talk about here is similar to convergence for ordinary sequences, but for random variables. Therefore, probabilities will play an important role. First, we ignore probabilities.

Definition 125 (Big-Oh)

Let $\{a_n\}_{n=1}^{\infty}$ be a sequence of constants. We say that $a_n = \mathcal{O}(1)$ if there exists a constant C such that

$$|a_n| \leq C$$

for all $n \geq 1$.

Remark 126. This is saying that the sequence $\{a_n\}_{n=1}^{\infty}$ is bounded by the constant C .

Generally, if we have two sequences of constants $\{a_n\}_{n=1}^{\infty}$ and $\{b_n\}_{n=1}^{\infty}$, we want to be able to relate them to each other. The following definition shows us how:

Definition 127 (Big-Oh Part Two)

Let $\{a_n\}_{n=1}^{\infty}$ and $\{b_n\}_{n=1}^{\infty}$ be sequences of constants. Then, we can write that $a_n = \mathcal{O}(b_n)$ if

$$\frac{a_n}{b_n} = \mathcal{O}(1)$$

Remark 128. This is equivalent to saying that the sequence of ratios is bounded

$$\left| \frac{a_n}{b_n} \right| \leq C$$

for all $n \geq 1$. We can see that $\mathcal{O}(1)$ is a special case of $\mathcal{O}(b_n)$ whereby $b_n = 1$.

Definition 129 (Little-Oh)

Let $\{a_n\}_{n=1}^{\infty}$ be a sequence of constants. Then, suppose that:

$$a_n \rightarrow 0$$

We can then write that $a_n = o(1)$.

Remark 130. Note that this is referring to ordinary convergence of sequences going to 0.

Example 131

Let $a_n = \frac{1}{n}$. Clearly

$$\frac{1}{n} \rightarrow 0$$

as $n \rightarrow \infty$. Then, we can write that $a_n = o(1)$.

Definition 132 (Little-Oh Part Two)

Let $\{a_n\}_{n=1}^{\infty}$ and $\{b_n\}_{n=1}^{\infty}$ be sequences of constants. Then, suppose that:

$$\frac{a_n}{b_n} \rightarrow 0$$

We can then write that $a_n = o(b_n)$.

Remark 133. If you compare this definition to our first definition of $o(1)$, $a_n = o(1)$ is the case of $a_n = o(b_n)$ where $b_n = 1$.

We now want to talk about what it means for a random variable to be **bounded**. Since it is a random variable, it does not make sense to talk about being bounded without involving probabilities. Therefore, we can show how to define what it means for a random variable being stochastically bounded.

Definition 134 (Bounded in probability)

Let $\{X_n\}_{n=1}^{\infty}$ be a sequence of random variables. We say that the sequence of random variables X_n is **bounded in probability** if for every $\epsilon > 0$, there exists constants δ and N such that

$$\mathbb{P}\left[|X_n| > \delta\right] < \epsilon$$

for all $n > N$. We write that

$$X_n = O_p(1) \tag{59}$$

Like what we did before, we can also extend our definition of being bounded in probability.

Definition 135 (Bounded in probability Part Two)

Let $\{X_n\}_{n=1}^{\infty}$ be a sequence of random variables and let $\{a_n\}_{n=1}^{\infty}$ be a sequence of constants. We say that the sequence of random variables is **bounded in probability** if for every $\epsilon > 0$, there exists constants δ and N such that

$$\mathbb{P}\left[\left|\frac{X_n}{a_n}\right| > \delta\right] < \epsilon$$

for all $n > N$. We write that

$$X_n = O_p(a_n) \tag{60}$$

Let us break down what this is saying. We are saying that if we take our sequence of random variables X_n , divide it by our sequence of constants a_n , then the probability that this ratio X_n/a_n is greater than δ is smaller than ϵ (which we defined to be a small number) for all random variables X_n where their index $n > N$. N here could be any number such as 100. If that was the case, then that means $X_{101}, X_{102}, X_{103}, \dots$ will all satisfy this property:

$$\mathbb{P}\left[\left|\frac{X_n}{a_n}\right| > \delta\right] < \epsilon$$

for $n \geq 101$.

We can also define an equivalent to little o in terms of probability.

Definition 136 (Little op)

Let $\{X_n\}_{n=1}^{\infty}$ be a sequence of random variables. We say that X_n converges in probability to zero if for every $\epsilon > 0$,

$$\lim_{n \rightarrow \infty} \mathbb{P}\left(|X_n| \geq \epsilon\right) = 0 \tag{61}$$

That is,

$$X_n \xrightarrow{P} 0$$

We write that

$$X_n = o_p(1) \tag{62}$$

Definition 137 (Little op Part Two)

Let $\{X_n\}_{n=1}^{\infty}$ be a sequence of random variables and let $\{a_n\}_{n=1}^{\infty}$ be a sequence of nonzero constants. We say that X_n/a_n converges in probability to zero if for every $\epsilon > 0$,

$$\lim_{n \rightarrow \infty} \mathbb{P}\left(\left|\frac{X_n}{a_n}\right| \geq \epsilon\right) = 0 \quad (63)$$

That is,

$$\frac{X_n}{a_n} \xrightarrow{P} 0$$

We write that

$$X_n = o_p(a_n) \quad (64)$$

When we talk about big-oh and little-oh notation, we are not actually talking about strict equality i.e. $X = o_p(1)$. It is instead a way for us to write the limits or talk about the convergence of random variables.

Example 138

Suppose that we have a sequence of random variables $\{X_n\}$ such that it converges in probability to X :

$$X_n \xrightarrow{P} X$$

Then, we can rewrite our expression for X_n as

$$X_n = X + o_p(1)$$

which is another way of writing that the difference $X_n - X$ goes to zero in probability.

We now show some interesting properties of $\mathcal{O}_p$ and o_p .

Lemma 139

$$o_p(1) + o_p(1) = o_p(1)$$

Example 140

Suppose that $X_n = o_p(1)$ and $Y_n = o_p(1)$. This implies that $X_n \xrightarrow{P} 0$ and $Y_n \xrightarrow{P} 0$. Then, by the continuous mapping theorem

$$X_n + Y_n \xrightarrow{P} 0$$

Lemma 141

$$\mathcal{O}_p(1) + o_p(1) = \mathcal{O}_p(1)$$

Proof. If one term is bounded in probability whilst the other term is going to 0, then the sum of the two will also be bounded in probability. $\square$

Lemma 142

$$\mathcal{O}_p(1)o_p(1) = o_p(1)$$

Proof. If one term is bounded in probability and the other term goes to 0 in probability, then their product will also go to 0 in probability. $\square$

Lemma 143

Let A be a nonzero constant. Then

$$o_p(A) = A \times o_p(1)$$

Proof. Suppose that $X_n = o_p(A)$. This implies that $\frac{X_n}{A} = o_p(1)$. Then, we get that

$$X_n = Ao_p(1).$$

$\square$

Theorem 144 (Convergence in distribution with big oh-pee)

If $X_n \xrightarrow{D} X$, then $X_n = \mathcal{O}_p(1)$.

Theorem 145 (Little oh-pee implies big oh-pee)

If $X_n = o_p(a_n)$, then $X_n = \mathcal{O}_p(a_n)$.

Remark 146. Little oh-pee is a stronger property than big oh-pee.

Proposition 147 (Linking deterministic and stochastic order)

Deterministic order implies stochastic order:

1. If $X_n = o(1)$ then $X_n = o_p(1)$
2. If $X_n = \mathcal{O}(1)$ then $X_n = \mathcal{O}_p(1)$

Example 148

Suppose that we have that $n^{-1} \rightarrow 0$ and $X_n \xrightarrow{D} X$. First, we have that $X_n = \mathcal{O}_p(1)$. We also have that $n^{-1} = o(1) = o_p(1)$. Therefore, we can conclude that:

$$n^{-1}X_n = o_p(1)\mathcal{O}_p(1) = o_p(1).$$

6 Maximum Likelihood Estimation

6.1 Introduction and Set Up

Suppose we had observed data $\mathbf{y} = (y_1, \dots, y_n)$. We assume that this data has the **joint probability density**:

$$f(\mathbf{y}; \boldsymbol{\theta}) = f(y_1, \dots, y_n; \boldsymbol{\theta})$$

We assume that the model is **correctly specified**. This means that there exists a parameter $\boldsymbol{\theta}_0 \in \Theta$ such that the data is actually generated by the density function:

$$\mathbf{y} \sim f(\cdot; \boldsymbol{\theta}_0).$$

Here, the data is **fixed** whilst the parameter $\boldsymbol{\theta}$ varies.

Definition 149 (Likelihood Function)

The **likelihood** function for $\mathbf{y}$ is defined as:

$$\mathcal{L}(\boldsymbol{\theta}; \mathbf{y}) \equiv f(\mathbf{y}; \boldsymbol{\theta}) \quad (65)$$

where $f(\mathbf{y}; \boldsymbol{\theta})$ is the joint density function of $\mathbf{y}$. The **log-likelihood** function is the logarithm of the likelihood function:

$$\ell(\boldsymbol{\theta}; \mathbf{y}) \equiv \log(\mathcal{L}(\boldsymbol{\theta}; \mathbf{y})). \quad (66)$$

In general, it is easier to compute the likelihood function if the observations are i.i.d.

Proposition 150 (Likelihood Function Under I.I.D Data)

Suppose that the observations $y_1, \dots, y_n$ are i.i.d. Then, the likelihood function for $\mathbf{y}$ is:

$$\mathcal{L}(\boldsymbol{\theta}; \mathbf{y}) = \prod_{i=1}^n f(y_i; \boldsymbol{\theta}). \quad (67)$$

The log-likelihood function is given by:

$$\ell(\boldsymbol{\theta}; \mathbf{y}) = \log(\mathcal{L}(\boldsymbol{\theta}; \mathbf{y})) = \sum_{i=1}^n \log(f(y_i; \boldsymbol{\theta})) \quad (68)$$

Remark 151. Suppose that we fixed the parameter $\boldsymbol{\theta}$. Now, a very important thing to realise is that $\ell(\boldsymbol{\theta}; \mathbf{y})$ is a **random variable** as it is just a function of the random variables $\mathbf{y}$ when we fixed $\boldsymbol{\theta}$.

The likelihood function measures the relative support that the observed data $\mathbf{y}$ provide for different values of the parameter $\boldsymbol{\theta}$. Given this, we want to find the parameter value with the greatest likelihood. This leads us to the maximum likelihood estimator.

Definition 152 (Maximum Likelihood Estimator)

The **maximum likelihood estimator** is the estimator $\hat{\theta}_{MLE}$ of the parameter θ that maximises the (log) likelihood function given the data $\mathbf{y}$. That is,

$$\hat{\theta}_{MLE} = \arg \max_{\theta \in \Theta} \mathcal{L}(\theta; \mathbf{y}) = \arg \max_{\theta \in \Theta} \ell(\theta; \mathbf{y})$$

Remark 153. Notice that due to monotonicity, maximising the likelihood function or the log-likelihood function gives us the same result for the MLE.

The maximum likelihood estimator $\hat{\theta}_{MLE}$ maximises the log-likelihood function such that for any other estimator $\tilde{\theta}$,

$$\ell(\hat{\theta}_{MLE}; \mathbf{y}) \geq \ell(\tilde{\theta}; \mathbf{y})$$

for all possible parameters $\theta \in \Theta$.

6.2 Score Function and Fisher Information Matrix

We have defined the log-likelihood function and the MLE as the estimate that maximises this function. Hence, MLE is the process of finding the maximum of the log-likelihood function. Now, the sharper the curvature of the log-likelihood function around the maximum point, the more accurate it will be to locate this maximum point. The more curved a log-likelihood function is, the more accurate our MLE will be. Hence, we are interested in analysing the derivatives of the log-likelihood function. These derivatives have important names in Statistics.

Definition 154 (Score Function)

The random variable that is the first order derivatives of the log-likelihood function is called the **Score function**:

$$\frac{\partial \ell(\theta; \mathbf{y})}{\partial \theta} \tag{69}$$

The Score function indicates the steepness of the log-likelihood function and thereby the sensitivity to infinitesimal changes to the parameter values.

Example 155

Suppose that $Y_i \sim \mathcal{N}(\mu, \sigma^2)$. Then, the score function is given by:

$$\frac{\partial \ell(\theta; \mathbf{y})}{\partial \theta} = \left(\frac{\partial \ell}{\partial \mu}, \frac{\partial \ell}{\partial \sigma^2} \right)^T.$$

We can in fact characterise the MLE by the Score function.

Proposition 156 (Necessary Conditions for MLE)

Suppose that the log-likelihood function $\ell(\theta; \mathbf{y})$ is differentiable and concave and that the MLE is an interior point of Θ . Then, the necessary condition for the MLE is that the Score function is equal to zero:

$$\frac{\partial \ell(\theta; \mathbf{y})}{\partial \theta} = \mathbf{0}.$$

Proof. This immediately follows from the maximisation program for the MLE:

$$\hat{\theta}_{MLE} = \arg \max_{\theta \in \Theta} \mathcal{L}(\theta; \mathbf{y}) = \arg \max_{\theta \in \Theta} \ell(\theta; \mathbf{y}).$$

□

As the **Score function is a random variable**, we want to compute the expectation and the variance of the this random variable.

Proposition 157 (Expectation of Score Function)

Suppose that the support of $\mathbf{y}$ does not depend on θ and that differentiation and integration can be interchanged. Then, the expectation of the score function is zero:

$$\mathbb{E}_{\mathbf{y}} \left[\frac{\partial \ell(\theta)}{\partial \theta} \right] = \mathbf{0} \quad (70)$$

Proof. First, recall that by definition of a probability density function:

$$\int_{\mathcal{Y}} f(\mathbf{y}; \theta) d\mathbf{y} = 1$$

where $\mathcal{Y}$ is the support of the random variable $\mathbf{y}$. Now, differentiate both sides with respect to θ :

$$\int_{\mathcal{Y}} \frac{\partial f(\mathbf{y}; \theta)}{\partial \theta} d\mathbf{y} = \mathbf{0}$$

Now, divide and multiply by $f(\mathbf{y}; \theta)$ on the left hand side:

$$\int_{\mathcal{Y}} \frac{\frac{\partial f(\mathbf{y}; \theta)}{\partial \theta}}{f(\mathbf{y}; \theta)} f(\mathbf{y}; \theta) d\mathbf{y} = \mathbf{0} \quad (71)$$

Now notice that in Eq. (71), the first term in the integral is actually the derivative of the logarithm of the density function by the chain rule:

$$\frac{\frac{\partial f(\mathbf{y}; \theta)}{\partial \theta}}{f(\mathbf{y}; \theta)} = \frac{\partial \ln(f(\mathbf{y}; \theta))}{\partial \theta} =_{(1)} \frac{\partial \ell(\theta)}{\partial \theta}$$

whereby $=_{(1)}$ is the realisation that the derivative of the log density is the definition of the score function! We can then plug this into Eq. (71) to get:

$$\int_{\mathcal{Y}} \frac{\partial \ell(\theta)}{\partial \theta} f(\mathbf{y}; \theta) d\mathbf{y} = \mathbf{0}$$

However, this by definition is just the expectation of the score function with respect to the density of the random variable $\mathbf{y}$, so we can replace the left hand side with the expectation of the score function and we are done:

$$\mathbb{E}_{\mathbf{y}} \left[\frac{\partial \ell(\theta)}{\partial \theta} \right] = \mathbf{0}$$

□

Proposition 158 (Variance of Score Function)

Under the same regularity conditions as above, the **variance of the score vector** is the negative of the expectation of the Hessian of the log-likelihood function:

$$Var\left[\frac{\partial \ell(\theta)}{\partial \theta}\right] = -\mathbb{E}\left[\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^T}\right] \quad (72)$$

Proof. First, we use the fact that the Score function has an expectation of zero:

$$\mathbb{E}_y\left[\frac{\partial \ell(\theta)}{\partial \theta}\right] = \int_y \frac{\partial \ell(\theta)}{\partial \theta} f(y; \theta) dy = \mathbf{0}$$

Now, we differentiate with respect to θ (and assume we can interchange differentiation and integration):

$$\begin{aligned} \int_y \frac{\partial}{\partial \theta^T} \left[\frac{\partial \ell(\theta)}{\partial \theta} f(y; \theta) \right] dy &= \mathbf{0} \\ \Leftrightarrow \int_y \left[\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^T} f(y; \theta) + \frac{\partial \ell(\theta)}{\partial \theta} \frac{\partial f(y; \theta)}{\partial \theta^T} \right] dy &= \mathbf{0} \\ \Leftrightarrow \int_y \left[\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^T} + \frac{\partial \ell(\theta)}{\partial \theta} \frac{\frac{\partial f(y; \theta)}{\partial \theta^T}}{f(y; \theta)} \right] f(y; \theta) dy &= \mathbf{0} \\ \Leftrightarrow \int_y \left[\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^T} + \frac{\partial \ell(\theta)}{\partial \theta} \frac{\partial \ell(\theta)}{\partial \theta^T} \right] f(y; \theta) dy &= \mathbf{0} \\ \Leftrightarrow \int_y \frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^T} f(y; \theta) dy + \int_y \frac{\partial \ell(\theta)}{\partial \theta} \frac{\partial \ell(\theta)}{\partial \theta^T} f(y; \theta) dy &= \mathbf{0} \end{aligned}$$

Both integrals are just expectations with respect to y :

$$\mathbb{E}_y\left[\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^T}\right] + \mathbb{E}_y\left[\left(\frac{\partial \ell(\theta)}{\partial \theta}\right)\left(\frac{\partial \ell(\theta)}{\partial \theta}\right)^T\right] = \mathbf{0}$$

So then we have that the second moment of the score function is given by the negative expectation of the Hessian of the log-likelihood function:

$$\mathbb{E}_y\left[\left(\frac{\partial \ell(\theta)}{\partial \theta}\right)\left(\frac{\partial \ell(\theta)}{\partial \theta}\right)^T\right] = -\mathbb{E}_y\left[\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^T}\right]$$

Now, we know that the score function has mean zero, this then implies for the variance of the score function:

$$Var\left[\frac{\partial \ell(\theta)}{\partial \theta}\right] = \mathbb{E}_y\left[\left(\frac{\partial \ell(\theta)}{\partial \theta}\right)\left(\frac{\partial \ell(\theta)}{\partial \theta}\right)^T\right] - \mathbb{E}\left[\frac{\partial \ell(\theta)}{\partial \theta}\right]\mathbb{E}\left[\frac{\partial \ell(\theta)}{\partial \theta}\right]^T = \mathbb{E}_y\left[\left(\frac{\partial \ell(\theta)}{\partial \theta}\right)\left(\frac{\partial \ell(\theta)}{\partial \theta}\right)^T\right] - \mathbf{0}$$

We can then plug in our expression for the second moment to get our desired result:

$$Var\left[\frac{\partial \ell(\theta)}{\partial \theta}\right] = -\mathbb{E}_y\left[\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^T}\right]$$

□

The variance of the Score function actually has an important name in Statistics.

Definition 159 (Fisher Information Matrix)

The **Fisher Information matrix** is the variance of the Score function:

$$I(\theta) = -\mathbb{E} \left[\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^T} \right] \quad (73)$$

Notice, that the **Fisher information matrix** is the negative of the expectation of the **Hessian** of the log-likelihood function $\ell(\theta)$. This is also equivalent to saying that it is the negative of the expectation of the **gradient** of the Score function (which itself is the derivative of the log-likelihood function):

$$\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^T} = \frac{\partial}{\partial \theta^T} \left[\frac{\partial \ell(\theta)}{\partial \theta} \right]$$

Now, this derivative of the Score function **without taking the expectation** also has an important name in Statistics.

Definition 160 (Observed Fisher Information Matrix)

The **Observed Fisher Information matrix** is the gradient of the Score function:

$$I_o(\theta) = -\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^T} \quad (74)$$

Remark 161. Notice that the Fisher Information matrix and the Observed Fisher Information matrix are very similar objects. The only difference is that the Fisher Information matrix has an expectation operator whilst the Observed Fisher Information matrix does not. In practice, we would compute the Observed Fisher information matrix by computing the Hessian of the log-likelihood function. We would then take the expectation of the Observed Fisher information matrix to retrieve the Fisher Information matrix.

Example 162

(Geometric Distribution). Suppose we had $\mathbf{X} = (X_1, \dots, X_n)$ i.i.d random variables whereby $X_i \sim \text{Geo}(\pi)$ with the probability density function:

$$Pr(X_i = x_i) = (1 - \pi)^{x_i} \pi$$

for $x_i = 0, 1, \dots$

The **log-likelihood function** is given by:

$$\ell(\pi; \mathbf{X}) = \sum_{i=1}^n \log(Pr(X_i = x_i)) = \sum_{i=1}^n \left[x_i \log(1 - \pi) + \log(\pi) \right] = n[\bar{x} \log(1 - \pi) + \log(\pi)]$$

whereby $\bar{x}$ is the sample mean.

The **Score function** is given by:

$$\frac{d\ell(\pi; \mathbf{X})}{d\pi} = n \left(\frac{1}{\pi} - \frac{\bar{x}}{1 - \pi} \right)$$

We can solve for the **maximum likelihood estimator** by setting the Score function to zero and solving for π to get:

$$\begin{aligned} n \left(\frac{1}{\hat{\pi}_{MLE}} - \frac{\bar{x}}{1 - \hat{\pi}_{MLE}} \right) &= 0 \\ \Rightarrow \hat{\pi}_{MLE} &= \frac{1}{1 + \bar{x}} \end{aligned}$$

The **Observed Fisher Information matrix** is the negative of the derivative of the Score function:

$$I_o = -\frac{d}{d\pi} \left[\frac{d\ell(\pi; \mathbf{X})}{d\pi} \right] = -\frac{d}{d\pi} \left[n \left(\frac{1}{\pi} - \frac{\bar{x}}{1 - \pi} \right) \right] = n \left[\frac{1}{\pi^2} + \frac{\bar{x}}{(1 - \pi)^2} \right]$$

We now want to compute the **Fisher Information matrix**. To do so, we take the expectation of the Fisher Information matrix (remember that we need only x is the random variable here):

$$I = \mathbb{E}[I_o] = \mathbb{E} \left[n \left[\frac{1}{\pi^2} + \frac{\bar{x}}{(1 - \pi)^2} \right] \right] = \frac{n}{\pi^2} + \frac{n\mathbb{E}[\bar{x}]}{(1 - \pi)^2}$$

Now, the expectation of a Geometric random variable is $\mathbb{E}[X_i] = (1 - \pi)/\pi$. This implies that the expectation of the sample average is:

$$\mathbb{E}[\bar{x}] = \frac{1}{n} \sum_{i=1}^n \mathbb{E}[X_i] = \frac{1}{n} \left(n \frac{1 - \pi}{\pi} \right) = \frac{1 - \pi}{\pi}$$

Therefore, we plug this in to get our final expression for the Fisher information matrix:

$$I = \frac{n}{\pi^2} + \frac{n\mathbb{E}[\bar{x}]}{(1 - \pi)^2} = \frac{n}{\pi^2} + \frac{n \frac{1 - \pi}{\pi}}{(1 - \pi)^2} = \frac{n}{\pi^2(1 - \pi)}.$$

We are usually interested in comparing the performance of estimators. In particular, we usually compare the efficiency between estimators. However, for technical reasons, we usually require to restrict our attention to the class of estimators that are **unbiased**. Within this class of estimators, we can actually derive what is the **lowest variance that any unbiased estimator** can achieve. This is known as the **Cramer-Rao lower bound** and it turns out that in fact, it is the inverse of the **Fisher Information matrix**.

Theorem 163 (Cramer–Rao Lower Bound)

Let $\tilde{\theta}(\cdot)$ be an **unbiased estimator** and suppose that the usual regularity conditions hold. Then the covariance matrix of any unbiased estimator is bounded below, in the positive-semidefinite ordering, by the inverse of the Fisher Information matrix:

$$\text{Var}[\tilde{\theta}] \succeq \mathbf{I}(\theta)^{-1}. \quad (75)$$

Therefore, if we have an unbiased estimator $\hat{\theta}$ and it has a variance that is equal to the inverse of the Fisher Information matrix, then this estimator has achieved the lowest possible variance any unbiased estimator can achieve and is hence efficient.

6.3 MLE for Regression

Suppose that we had the data $(y_i, \mathbf{x}_i)$ where y_i is the outcome and $\mathbf{x}_i$ are the covariates. Then, a key assumption is that we can factor the joint density as:

$$p(\mathbf{y}, \mathbf{x}; \theta) = q(\mathbf{y}|\mathbf{x}; \theta)r(\mathbf{x}) \quad (76)$$

whereby the joint density can be written as a product of the conditional density and the marginal density. Notice that q depends on the parameter θ but the marginal density r does not.

Under the assumption of i.i.d data, this means we can write the **conditional log-likelihood** as:

$$\ell_n(\theta) = \sum_{i=1}^n \log(p(y_i, \mathbf{x}_i; \theta)) = \sum_{i=1}^n \log(q(y_i|\mathbf{x}_i; \theta)) + \sum_{i=1}^n \log(r(\mathbf{x}_i))$$

However, notice that when we are maximising over θ , the second term does not contain θ and hence we can drop it. Therefore, the conditional log-likelihood is given by:

$$\ell_n(\theta) = \sum_{i=1}^n \log(q(y_i|\mathbf{x}_i; \theta)) \quad (77)$$

You can then maximise this to find your estimator.

6.3.1 Homoskedastic Gaussian Regression

We show how to derive the MLE for a regression model with Gaussian error terms. Suppose we had the linear regression model:

$$y_i|\mathbf{x}_i \sim \mathcal{N}(\mathbf{x}_i^T \boldsymbol{\beta}, \sigma^2) \quad (78)$$

which is equivalent to writing:

$$y_i = \mathbf{x}_i^T \boldsymbol{\beta} + \sigma u_i$$

where $u_i \sim \mathcal{N}(0, 1)$ and is independent of $\mathbf{x}_i$. Here, we want to estimate both $\boldsymbol{\beta}$ and σ . Now, the joint density function $p(\cdot)$ is given by the conditional density $q(\cdot)$ times the marginal density $r(\cdot)$:

$$p(\mathbf{y}, \mathbf{X}; \boldsymbol{\beta}, \sigma) = q(\mathbf{y}|\mathbf{X}; \boldsymbol{\beta}, \sigma) \cdot r(\mathbf{x}).$$

Now, when we take logs, we can see that the marginal density $r(\mathbf{x})$ does not depend on $\boldsymbol{\beta}$ or σ so we can ignore

it. Then, we can then write the conditional log-likelihood as:

$$\begin{aligned}\ell_n^c(\boldsymbol{\beta}, \sigma) &= \sum_{i=1}^n \log\{q(y_i | \mathbf{x}_i; \boldsymbol{\beta}, \sigma)\} \\ &= \sum_{i=1}^n \log \left\{ \frac{1}{\sigma} \phi \left(\frac{y_i - \mathbf{x}_i^T \boldsymbol{\beta}}{\sigma} \right) \right\} \\ &= -\frac{n}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} \sum_{i=1}^n (y_i - \mathbf{x}_i^T \boldsymbol{\beta})^2,\end{aligned}$$

Now, if we wanted to find the estimator for $\boldsymbol{\beta}$, we can see that only the second term depends on $\boldsymbol{\beta}$. Therefore, we need to find a maximum likelihood estimator $\hat{\boldsymbol{\beta}}$ that **minimises** the sum of squares:

$$\sum_{i=1}^n (y_i - \mathbf{x}_i^T \boldsymbol{\beta})^2.$$

However, this is just the definition of the OLS estimator! This gives us a useful result.

Theorem 164 (OLS and MLE Coincide)

The OLS estimator and the MLE estimator coincide for $\boldsymbol{\beta}$ in the linear regression model with Gaussian error terms.

Now, it can be shown that taking first order conditions, we can derive the MLE $\hat{\sigma}^2$ to be:

$$\hat{\sigma}_{MLE}^2 = \frac{1}{n} \sum_{i=1}^n (y_i - \mathbf{x}_i^T \hat{\boldsymbol{\beta}})^2 \quad (79)$$

This differs from the unbiased OLS variance estimator, which divides the residual sum of squares by $n - K$.

We now show a similar argument using a multivariate normal distribution.

Definition 165 (Multivariate Normal Distribution)

The $n \times 1$ random vector $\mathbf{Z}$ has the multivariate normal distribution $\mathbf{Z} \sim \mathcal{N}(\boldsymbol{\mu}, \Sigma)$ if its probability density function is:

$$f_{\mathbf{Z}}(\mathbf{Z} = \mathbf{a}) = (2\pi)^{-\frac{n}{2}} \frac{1}{\sqrt{\det(\Sigma)}} \exp \left(-\frac{1}{2} (\mathbf{a} - \boldsymbol{\mu})^T \Sigma^{-1} (\mathbf{a} - \boldsymbol{\mu}) \right). \quad (80)$$

Now, suppose we had the classical linear regression model with stochastic regressors and conditional normality on the error terms:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u} \quad \text{where } \mathbf{u} | \mathbf{X} \sim \mathcal{N}(0, \sigma^2 \mathbf{I})$$

where $\mathbf{I}$ is the $n \times n$ identity matrix.

This is equivalent to writing that:

$$\mathbf{y} | \mathbf{X} \sim \mathcal{N}(\mathbf{X}\boldsymbol{\beta}, \sigma^2 \mathbf{I}) \quad (81)$$

Matching this to our notation for the multivariate normal, we can see that the mean and variance are:

$$\boldsymbol{\mu} = \mathbf{X}\boldsymbol{\beta} \quad \text{and} \quad \Sigma = \sigma^2 \mathbf{I}.$$

Therefore, the conditional density of $\mathbf{y}$ given $\mathbf{X}$ is:

$$f(\mathbf{y} | \mathbf{X}) = (2\pi\sigma^2)^{-n/2} \exp \left(-\frac{1}{2\sigma^2} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^T (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) \right) \quad (82)$$

We can then derive the conditional likelihood function by taking the summation over the sample:

$$\mathcal{L}(\boldsymbol{\beta}, \sigma^2) = (2\pi)^{-\frac{n}{2}} (\sigma^2)^{-\frac{n}{2}} \exp\left(\frac{-(\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^T (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})}{2\sigma^2}\right). \quad (83)$$

We can then take the logarithm to get the **log-likelihood** function:

$$\ell(\boldsymbol{\beta}, \sigma^2) = -\frac{n}{2} \log(2\pi) - \frac{n}{2} \log(\sigma^2) - \frac{1}{2\sigma^2} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^T (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) \quad (84)$$

Again, notice that only the final term depends on $\boldsymbol{\beta}$. Since the quadratic enters with a negative sign, to find the MLE for $\boldsymbol{\beta}$, we minimise:

$$(\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^T (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) \quad (85)$$

which is equivalent to minimising the sum of squared errors with respect to $\boldsymbol{\beta}$. Therefore, again we find that the conditional maximum likelihood estimator $\hat{\boldsymbol{\beta}}_{MLE}$ in the classical linear regression model with normally distributed errors is identical to the OLS estimator $\hat{\boldsymbol{\beta}}_{OLS}$. Similarly, we can derive the MLE of the variance to be:

$$\hat{\sigma}_{MLE}^2 = \frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{n} \quad (86)$$

6.4 Asymptotic Theory

The MLE has some really nice properties.

Theorem 166 (Consistency of MLE)

Under standard regularity and identification conditions, the MLE is a consistent estimator

$$\hat{\boldsymbol{\theta}}_{MLE} \xrightarrow{P} \boldsymbol{\theta}_0 \quad (87)$$

where $\boldsymbol{\theta}_0$ is the true parameter value.

Proof. First, for each value $\boldsymbol{\theta} \in \Theta$, we have that the average log-likelihood is given by:

$$\bar{\ell}_n(\boldsymbol{\theta}) = \frac{1}{n} \sum_{i=1}^n \log(p(\mathbf{x}_i; \boldsymbol{\theta}))$$

Now, notice that as $\boldsymbol{\theta}$ is fixed, this means that $\log(p(\mathbf{x}_i; \boldsymbol{\theta}))$ are i.i.d random variables. Therefore, we can apply the weak law of large numbers:

$$\frac{1}{n} \sum_{i=1}^n \log(p(\mathbf{x}_i; \boldsymbol{\theta})) \xrightarrow{P} \mathbb{E}_{\boldsymbol{\theta}_0}[\log(p(\mathbf{x}_i; \boldsymbol{\theta}))] = \int_{\mathcal{X}} \log(p(w; \boldsymbol{\theta})) p(w; \boldsymbol{\theta}_0) dw$$

where notice that the expectation is taken with respect to the true population value $\boldsymbol{\theta}_0$.

Heuristically, we can say that the average log-likelihood $\bar{\ell}_n(\boldsymbol{\theta}) \rightarrow \ell_0(\boldsymbol{\theta})$ which is the true log-likelihood.

Now, we want to argue heuristically:

$$\hat{\boldsymbol{\theta}}_{MLE} = \arg \max_{\boldsymbol{\theta} \in \Theta} \bar{\ell}_n(\boldsymbol{\theta}) \xrightarrow{P} \arg \max_{\boldsymbol{\theta} \in \Theta} \ell_0(\boldsymbol{\theta}) = \boldsymbol{\theta}_0 \quad (88)$$

or

$$\hat{\boldsymbol{\theta}}_{MLE} \xrightarrow{P} \boldsymbol{\theta}_0$$

where θ_0 is the true value. We hand wave our argument and say that since $\bar{\ell}_n(\theta) \rightarrow \ell_0(\theta)$ then as the functions:

$$\bar{\ell}_n(\hat{\theta}_{MLE}) \rightarrow \ell_0(\theta_0)$$

then the maximiser:

$$\hat{\theta}_{MLE} \xrightarrow{P} \theta_0.$$

Note that to rigorously make this argument, we need to show uniform convergence, not pointwise convergence. $\square$

We also have the distribution of the MLE.

Theorem 167 (MLE is Asymptotically Normal)

Under standard regularity and identification conditions, the MLE is **asymptotically normal**:

$$n^{1/2}(\hat{\theta}_{MLE} - \theta_0) \xrightarrow{D} \mathcal{N}(\mathbf{0}, \mathbf{H}^{-1} \mathbf{S} \mathbf{H}^{-1}) \quad (89)$$

where $\mathbf{H} = -\mathbb{E}[\partial^2 \ell_i(\theta_0) / (\partial \theta \partial \theta^T)]$, $\mathbf{S} = \mathbb{E}[\partial \ell_i(\theta_0) / \partial \theta \partial \ell_i(\theta_0) / \partial \theta^T]$ is the variance of the individual score, and θ_0 is the true parameter value.

Fortunately, we can simplify this expression for the asymptotic distribution of the MLE.

Theorem 168 (Efficiency of MLE)

Suppose that the model is correctly specified and the standard regularity conditions hold. Then, the information matrix equality gives $\mathbf{H} = \mathbf{S}$ and the MLE attains the Cramer–Rao lower bound:

$$n^{1/2}(\hat{\theta}_{MLE} - \theta_0) \xrightarrow{D} \mathcal{N}(\mathbf{0}, \mathbf{H}^{-1}) \quad (90)$$

6.5 Hypothesis Testing

We now look at three different types of tests.

6.5.1 Likelihood Ratio Test

The intuition for the likelihood ratio test is as follows. Suppose we have an unrestricted model where the parameter space is Θ . From this, we can compute the maximum likelihood of the model under this space of parameters:

$$\max_{\theta \in \Theta} \mathcal{L}(\theta).$$

Now, suppose we impose a restriction R which shrinks the parameter space to $\Theta_R \subset \Theta$. Now, as this is a smaller set, the maximum likelihood over this space should be smaller:

$$\max_{\theta \in \Theta_R} \mathcal{L}(\theta).$$

The question is how big of a drop in the likelihood this restriction creates. The bigger the drop, the more likely we are to reject the null hypothesis. We can also show that the resulting test statistic is asymptotically χ^2 .

Definition 169 (Likelihood Ratio Test)

Let Θ be the unrestricted parameter space. Let Θ_R be the restricted parameter space under the hypothesis H_R . Then, the likelihood-ratio statistic is defined as:

$$LR = -2 \log \left[\frac{\max_{\theta \in \Theta_R} \mathcal{L}(\theta)}{\max_{\theta \in \Theta} \mathcal{L}(\theta)} \right] = -2 \left\{ \ell(\hat{\theta}_R) - \ell(\hat{\theta}) \right\} \xrightarrow{D} \chi^2(p) \quad (91)$$

under the null hypothesis, where $\hat{\theta}$ and $\hat{\theta}_R$ are the unrestricted and restricted maximum likelihood estimators and $p = \dim \Theta - \dim \Theta_R$.

6.5.2 Wald Test

The Wald test asks the question on how big is the distance between the unrestricted estimator $\hat{\theta}$ and the restricted estimator $\hat{\theta}_R$ in the parameter space $\theta \in \Theta = \mathbb{R}^k$ under the null hypothesis:

$$H_0 : R^T \theta = r$$

for some known restriction $R \in \mathbb{R}^{k \times p}$ and $r \in \mathbb{R}^p$.

Definition 170 (Wald Test Statistic)

The Wald test statistic is given by:

$$W = (R^T \hat{\theta} - r)^T \left[R^T [I(\hat{\theta})]^{-1} R \right]^{-1} (R^T \hat{\theta} - r) \xrightarrow{D} \chi^2(p) \quad (92)$$

where I is the Fisher information matrix.

6.5.3 Lagrange Multiplier Test

First, recall that the score function is zero when evaluated the maximum likelihood estimator. In this case, it is zero when evaluated at the **unrestricted** maximum likelihood estimator $\hat{\theta}$. The Lagrange Multiplier test looks at the value of the score at the restricted estimator $\hat{\theta}_R$, knowing that this won't be zero and compares it to the score function evaluated at the unrestricted MLE. The test statistic associated is also asymptotically χ^2 .

Definition 171 (LM Test)

The Lagrange Multiplier test statistic is given by:

$$LM = [\dot{\ell}(\hat{\theta}_R)]^T \left[I(\hat{\theta}_R) \right]^{-1} [\dot{\ell}(\hat{\theta}_R)] \xrightarrow{D} \chi^2(p) \quad (93)$$

where I is the Fisher information matrix and $\dot{\ell}(\theta)$ is the Score function.

7 The General Regression Model

7.1 Sample and Population Models

This section recaps the things we have covered so far and re-establish notation for the following sections. We are interested in single equation linear models of the form:

$$y_i = \beta_1 x_{1i} + \beta_2 x_{2i} + \dots + \beta_K x_{Ki} + u_i \quad (94)$$

whereby for each observation i , we have a scalar outcome variable y_i , an error term u_i , and a set of K explanatory variables $x_{1i}, x_{2i}, \dots, x_{Ki}$. We then have a sample of n observations indexed by $i = 1, 2, \dots, n$. The linearity of this model refers to the fact that the parameters $\beta_1, \dots, \beta_K$ and the error term u_i enter the model additively.

We now denote $\mathbf{x}_i$ and $\boldsymbol{\beta}$ to be the $K \times 1$ column vectors

$$\mathbf{x}_i = \begin{bmatrix} x_{1i} \\ x_{2i} \\ \cdot \\ \cdot \\ \cdot \\ x_{Ki} \end{bmatrix} \quad \text{and} \quad \boldsymbol{\beta} = \begin{bmatrix} \beta_1 \\ \beta_2 \\ \cdot \\ \cdot \\ \cdot \\ \beta_K \end{bmatrix}$$

We can therefore rewrite Eq. (94) as:

$$y_i = \mathbf{x}_i^T \boldsymbol{\beta} + u_i \quad \text{for } i = 1, \dots, n. \quad (95)$$

In Eq. (95), this refers to the linear model for a **sample**. If we drop the subscripts, we get the **population model**:

$$y = \mathbf{x}^T \boldsymbol{\beta} + u \quad (96)$$

Definition 172 (Population Linear Model)

The matrix representation of the **population** version of the linear model is given by:

$$y = \mathbf{x}^T \boldsymbol{\beta} + u \quad (97)$$

Definition 173 (Sample Linear Model)

The matrix representation of the **sample** version of the linear model is given by:

$$y_i = \mathbf{x}_i^T \boldsymbol{\beta} + u_i \quad \text{for } i = 1, \dots, n. \quad (98)$$

Finally, we can stack the observations in Eq. (95) to have the vectors and matrices:

$$\mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix} \quad \mathbf{X} = \begin{bmatrix} \mathbf{x}_1^T \\ \mathbf{x}_2^T \\ \vdots \\ \mathbf{x}_n^T \end{bmatrix} = \begin{bmatrix} x_{11} & x_{21} & \cdot & \cdot & \cdot & x_{K1} \\ x_{12} & x_{22} & \cdot & \cdot & \cdot & x_{K2} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ x_{1n} & x_{2n} & \cdot & \cdot & \cdot & x_{Kn} \end{bmatrix} \quad \mathbf{u} = \begin{bmatrix} u_1 \\ u_2 \\ \vdots \\ u_n \end{bmatrix}$$

Here, $\mathbf{y}, \mathbf{u}$ are both $n \times 1$ vectors whilst $\mathbf{X}$ is a $n \times K$ matrix whereby each row represents all K observed regressors associated to a particular observation.

Definition 174 (Matrix Representation of Sample Linear Model)

The matrix representation of the sample version of the linear model is given by:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u} \quad (99)$$

We are interested in estimating the parameter vector $\boldsymbol{\beta}$ from the data on y_i and $x_{1i}, \dots, x_{Ki}$ for a sample of $i = 1, \dots, n$ observations. The **Ordinary Least Squares** (OLS) estimator finds the parameter vector $\boldsymbol{\beta}$ which minimizes the sum of the squared errors $u_i(\boldsymbol{\beta}) = y_i - \mathbf{x}_i^T \boldsymbol{\beta}$:

$$\hat{\boldsymbol{\beta}}_{OLS} = \arg \min_{\boldsymbol{\beta}} \sum_{i=1}^n u_i^2(\boldsymbol{\beta}) = \arg \min_{\boldsymbol{\beta}} \mathbf{u}(\boldsymbol{\beta})^T \mathbf{u}(\boldsymbol{\beta}) = \arg \min_{\boldsymbol{\beta}} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^T (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) \quad (100)$$

Definition 175 (Ordinary Least Squares Estimator)

The ordinary least squares estimator is the solution to the problem:

$$\hat{\boldsymbol{\beta}}_{OLS} = \arg \min_{\boldsymbol{\beta}} (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})^T (\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) \quad (101)$$

We can then take first order conditions of the ordinary least squares estimator problem and arrive at the system of K **normal equations** for the K unknown elements of $\boldsymbol{\beta}$:

$$\mathbf{X}^T \mathbf{y} = \mathbf{X}^T \mathbf{X} \boldsymbol{\beta} \quad (102)$$

If we know that the inverse $(\mathbf{X}^T \mathbf{X})^{-1}$ exist, then the parameter vector $\boldsymbol{\beta}$ which minimises the sum of squared residuals is given by:

$$\hat{\boldsymbol{\beta}}_{OLS} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} \quad (103)$$

This can be written in summation form:

$$\hat{\boldsymbol{\beta}}_{OLS} = \left(\sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^T \right)^{-1} \left(\sum_{i=1}^n \mathbf{x}_i y_i \right) \quad (104)$$

whereby we recall that $y_i = \mathbf{x}_i^T \boldsymbol{\beta} + u_i$ whereby $\mathbf{x}_i$ is a $K \times 1$ vector.

7.2 General Regression Model

We are interested in **predicting** the value taken by y_i using only the information taken by the random variables in $\mathbf{x}_i$. We can write our predictions using our regressors in the form: $h(\mathbf{x}_i)$ whereby $h(\mathbf{x}_i)$ is a function that maps from the K-dimensional vector $\mathbf{x}_i$ into the scalars. To measure how good our prediction, we can use the measure of mean squared error:

$$\mathbb{E}[(y_i - h(\mathbf{x}_i))^2].$$

The question is, what function $h(\cdot)$ can minimise the mean squared error of our prediction? The following theorem gives us the answer.

Theorem 176 (Unique Minimiser of Mean Squared Error)

The function $h^*(\mathbf{x}_i)$ that minimises the mean squared error:

$$MSE(h) = \mathbb{E}[\{y_i - h(\mathbf{x}_i)\}^2] \quad (105)$$

is **unique** and given by the **conditional expectation function**

$$h^*(\mathbf{x}_i) = \mathbb{E}[y_i | \mathbf{x}_i]. \quad (106)$$

This is why the conditional expectation function is really important for us. We now generalise our framework by introducing the general regression model. We then see that we can have two submodels arising from this: the linear regression model and the linear projection model.

Definition 177 (General Regression Model)

The general regression model with additive errors can be written as:

$$\mathbf{y} = \mathbb{E}[\mathbf{y} | \mathbf{X}] + \mathbf{u} \quad (107)$$

which is the additive decomposition of the random vector $\mathbf{y}$ into its conditional expectation given $\mathbf{X}$ and a vector of additive deviations $\mathbf{u}$ from this conditional expectation.

We get an interesting result regarding the conditional expectation.

Theorem 178 (Orthogonality Property of the Conditional Expectation Function)

The conditional expectation $h^*(\mathbf{x}_i) = \mathbb{E}[y_i | \mathbf{x}_i]$ has the following orthogonal property that for every other bounded function $g(\mathbf{x}_i)$:

$$\mathbb{E}\left[\left\{y_i - h^*(\mathbf{x}_i)\right\}g(\mathbf{x}_i)\right] = 0 \quad (108)$$

We get the following properties for the conditional expectation which follows from the orthogonality property of the conditional expectation function.

Corollary 179

We have the following properties for the conditional expectation function.

- If $\mathbf{x}_i$ is a constant, then $\mathbb{E}[y_i|\mathbf{x}_i] = \mathbb{E}[y_i]$;
- If y_i and $\mathbf{x}_i$ are independent, then $\mathbb{E}[y_i|\mathbf{x}_i] = \mathbb{E}[y_i]$;
- If $\mathbf{z}_i$ is independent of $(y_i, \mathbf{x}_i)$ then $\mathbb{E}[y_i|\mathbf{x}_i, \mathbf{z}_i] = \mathbb{E}[y_i|\mathbf{x}_i]$;

We can now define our first submodel.

Definition 180 (Linear Regression Model)

A linear regression model is obtained if the conditional expectation of $\mathbf{y}$ given $\mathbf{X}$ is a linear function of $\mathbf{X}$:

$$\mathbb{E}[\mathbf{y}|\mathbf{X}] = \mathbf{X}\boldsymbol{\beta}. \quad (109)$$

This implies that the general regression model:

$$\mathbf{y} = \mathbb{E}[\mathbf{y}|\mathbf{X}] + \mathbf{u}$$

becomes

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u} \quad (110)$$

The OLS estimator is really good in this linear regression model framework as it is able to estimate $\boldsymbol{\beta}$ under certain conditions. However, it is **not** the case that the conditional expectation is a linear function $\mathbb{E}[\mathbf{y}|\mathbf{X}] = \mathbf{X}\boldsymbol{\beta}$. In other words, we **do not necessarily** have that the true data generating process is:

$$\mathbf{y} = \mathbb{E}[\mathbf{y}|\mathbf{X}] + \mathbf{u} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}.$$

Whilst this may not be the true data generating process for our population data, we may still be interested in the linear projection of the population outcome variable $\mathbf{y}$ on the population regressors $\mathbf{X}$. The conditional expectation is the best unrestricted predictor under mean squared error, whereas the linear projection is the best predictor within the class of linear functions of $\mathbf{X}$.

Definition 181 (Linear Projection Model)

The linear projection model is when we **do not** assume that the conditional expectation is a linear function, but we are still interested on the **linear projection** of $\mathbf{y}$ on $\mathbf{X}$, which is denoted by

$$\mathbb{E}^*[\mathbf{y}|\mathbf{X}] = \mathbf{X}\boldsymbol{\beta}. \quad (111)$$

To summarise, we say that in our population model $\mathbf{y}$ and $\mathbf{X}$, we want to uncover the true relationship between the two. First, it is important to note that we can **always** project $\mathbf{y}$ onto $\mathbf{X}$ (this is just linear algebra). If we then assume that it is indeed the case that the relationship between the two variables is indeed that one is being projected onto the other, then we are under the linear regression model framework. This assumption of the outcome $\mathbf{y}$ being a projection onto the regressor $\mathbf{X}$ comes from the assumption that the conditional expectation is a linear function $\mathbb{E}[\mathbf{y}|\mathbf{X}] = \mathbf{X}\boldsymbol{\beta}$. If we do not make the assumption that the true relationship between the two variables is that $\mathbf{y}$ is a projection onto $\mathbf{X}$, we can still analyse this projection even though it is not the true relationship. This is the linear projection model.

7.3 Classical Linear Regression Models

We are now interested in establishing properties of the OLS estimators. However, in order to do so, we need to make more assumptions. We can first make assumptions that gives us properties that hold under finite samples. Alternatively, we can make assumptions that gives us properties that hold asymptotically.

Assumption 182 (Linear Conditional Expectation)

The conditional expectation of $\mathbf{y}$ given $\mathbf{X}$ is linear:

$$\mathbb{E}[\mathbf{y}|\mathbf{X}] = \mathbf{X}\boldsymbol{\beta} \quad (112)$$

$$\mathbb{E}[\mathbf{u}|\mathbf{X}] = 0 \quad (113)$$

Remark 183. *The assumption*

$$\mathbb{E}[\mathbf{u}|\mathbf{X}] = 0$$

is also known as **strict exogeneity**.

Assumption 184 (Conditional Homoskedasticity)

The conditional variance of $\mathbf{y}$ given $\mathbf{X}$ is:

$$\text{Var}[\mathbf{y}|\mathbf{X}] = \sigma^2 \mathbf{I} \quad (114)$$

$$\text{Var}[\mathbf{u}|\mathbf{X}] = \sigma^2 \mathbf{I} \quad (115)$$

Assumption 185 (Full Rank)

The matrix $\mathbf{X}$ has full rank:

$$\text{rank}(\mathbf{X}) = K \quad (116)$$

with probability one.

Remark 186. *The matrix $\mathbf{X}$ is a $n \times K$ matrix whereby we have that $K < n$. We know that the rank of a matrix is at most either the minimum of the column or row rank of a matrix.*

Assumption 187 (Independence of Observations)

We assume that observations $(y_i, \mathbf{x}_i)$ are independent over $i = 1, 2, \dots, n$, so that

$$\mathbb{E}[y_i | \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n] = \mathbb{E}[y_i | \mathbf{x}_i] = \mathbf{x}_i^T \boldsymbol{\beta} \quad (117)$$

or equivalently

$$\mathbb{E}[u_i | \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n] = \mathbb{E}[u_i | \mathbf{x}_i] = 0. \quad (118)$$

Independence also gives us:

Corollary 188

Independence of observations implies that:

$$Cov(y_i, y_j | \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n) = 0 \quad (119)$$

or equivalently:

$$Cov(u_i, u_j | \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n) = 0 \quad (120)$$

Corollary 189

Independence of observations together with the conditional-mean assumption implies that the error term u_i for observation i is uncorrelated with the explanatory variable $\mathbf{x}_j$ for all $i, j = 1, 2, \dots, n$.

7.4 Finite Properties in Regression Models

We now take the four assumptions stated in the previous subsection and show finite sample properties of the OLS estimator.

Proposition 190 (Conditional Unbiasedness of OLS Estimator)

The OLS estimator has the conditional expectation

$$\mathbb{E}[\hat{\beta}_{OLS}|\mathbf{X}] = \beta. \quad (121)$$

Proof.

$$\begin{aligned} \mathbb{E}[\hat{\beta}_{OLS}|\mathbf{X}] &= \mathbb{E}[(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}|\mathbf{X}] \\ &=_{(1)} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbb{E}[\mathbf{y}|\mathbf{X}] \\ &=_{(2)} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X} \beta \\ &= \beta \end{aligned}$$

whereby $=_{(1)}$ uses the fact that $\mathbf{X}$ is known when conditioned on $\mathbf{X}$ and $=_{(2)}$ uses the linear conditional expectation assumption $\mathbb{E}[\mathbf{y}|\mathbf{X}] = \mathbf{X}\beta$. $\square$

Theorem 191 (Unconditional Unbiasedness of OLS Estimator)

The OLS estimator is unbiased

$$\mathbb{E}[\hat{\beta}_{OLS}] = \beta. \quad (122)$$

Proof. We can apply the law of iterated expectations

$$\mathbb{E}[\hat{\beta}_{OLS}] =_{(1)} \mathbb{E}[\mathbb{E}[\hat{\beta}_{OLS}|\mathbf{X}]] = \mathbb{E}[\beta] = \beta$$

whereby we used the conditional expectation of the OLS estimator in $=_{(1)}$. $\square$

Theorem 192 (Conditional Variance of OLS Estimator)

$$\text{Var}[\hat{\beta}_{OLS}|\mathbf{X}] = \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1} \quad (123)$$

Proof.

$$\begin{aligned} \text{Var}[\hat{\beta}_{OLS}|\mathbf{X}] &= \text{Var}[(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}|\mathbf{X}] \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \text{Var}[\mathbf{y}|\mathbf{X}] ((\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T)^T \\ &=_{(1)} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \text{Var}[\mathbf{y}|\mathbf{X}] \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} \\ &=_{(2)} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \sigma^2 \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} \\ &= \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} \\ &= \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1} \end{aligned}$$

whereby we used the symmetry of $(\mathbf{X}^T \mathbf{X})^{-1}$ in $=_{(1)}$ and the assumption of homoskedasticity in $=_{(2)}$. $\square$

Theorem 193 (Unconditional Variance of OLS Estimator)

When $\mathbf{X}$ is stochastic, the unconditional variance of the OLS estimator is given by:

$$\text{Var}[\hat{\boldsymbol{\beta}}_{OLS}] = \sigma^2 \mathbb{E}[(\mathbf{X}^T \mathbf{X})^{-1}]. \quad (124)$$

Theorem 194 (Conditional Gauss-Markov Theorem)

The conditional Gauss-Markov theorem is given by

$$\text{Var}[\hat{\boldsymbol{\beta}}_{OLS}|\mathbf{X}] \leq \text{Var}[\tilde{\boldsymbol{\beta}}|\mathbf{X}] \quad (125)$$

for any other unbiased estimator $\tilde{\boldsymbol{\beta}}$ that is a linear function of the random vector $\mathbf{y}$ conditional on $\mathbf{X}$.

In practice, we don't know what σ^2 is but we would like to estimate it so that we can compute the conditional variance $\text{Var}[\hat{\boldsymbol{\beta}}_{OLS}|\mathbf{X}] = \sigma^2(\mathbf{X}^T \mathbf{X})^{-1}$. We can estimate it with the following estimator.

Proposition 195 (OLS Estimator of Variance)

The OLS estimator of the variance is given by:

$$\hat{\sigma}_{OLS}^2 = \frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{n - K} = \frac{\sum_{i=1}^n \hat{u}_i^2}{n - K} \quad (126)$$

whereby $\hat{\mathbf{u}} = \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}_{OLS}$ is the vector of OLS residuals and K is the number of regressors.

Proposition 196 (OLS Estimator of Variance is Unbiased)

The OLS estimator for the variance σ^2 is unbiased:

$$\mathbb{E}[\hat{\sigma}_{OLS}^2] = \sigma^2. \quad (127)$$

Proof. First, we note that the fitted residual can be expressed as:

$$\hat{\mathbf{u}} = \mathbf{M}\mathbf{y}$$

where $\mathbf{M}$ is the annihilator matrix. Proceeding further, we have that:

$$\begin{aligned} \hat{\mathbf{u}} &= \mathbf{M}\mathbf{y} \\ &= \mathbf{M}[\mathbf{X}\boldsymbol{\beta} + \mathbf{u}] \\ &= \mathbf{M}\mathbf{X}\boldsymbol{\beta} + \mathbf{M}\mathbf{u} \\ &\stackrel{(1)}{=} \mathbf{0} + \mathbf{M}\mathbf{u} \end{aligned}$$

whereby $\stackrel{(1)}{=}$ follows from the fact that $\mathbf{M}\mathbf{X} = \mathbf{0}$. Therefore, we have that

$$\hat{\mathbf{u}} = \mathbf{M}\mathbf{u}$$

We can therefore write the residual sum of squares as:

$$\hat{\mathbf{u}}^T \hat{\mathbf{u}} = \mathbf{u}^T \mathbf{M}^T \mathbf{M} \mathbf{u} = \mathbf{u}^T \mathbf{M} \mathbf{u},$$

where the last equality uses symmetry and idempotency of $\mathbf{M}$. Now, we look at the conditional expectation:

$$\begin{aligned}
\mathbb{E}[\hat{\mathbf{u}}^T \hat{\mathbf{u}} | \mathbf{X}] &= \mathbb{E}[\mathbf{u}^T \mathbf{M} \mathbf{u} | \mathbf{X}] \\
&=_{(1)} \mathbb{E}[\text{tr}(\mathbf{M} \mathbf{u} \mathbf{u}^T) | \mathbf{X}] \\
&=_{(2)} \text{tr}(\mathbf{M} \mathbb{E}[\mathbf{u} \mathbf{u}^T | \mathbf{X}]) \\
&= \text{tr}(\mathbf{M} \sigma^2 \mathbf{I}) \\
&= \sigma^2 \text{tr}(\mathbf{M}) \\
&=_{(4)} \sigma^2 (n - K)
\end{aligned}$$

whereby $=_{(1)}$ uses the trace representation of a quadratic form, $=_{(2)}$ uses the fact that $\mathbf{M}$ is known when conditioned on $\mathbf{X}$, and $=_{(4)}$ uses the fact that the trace of $\mathbf{M}$ is $(n - K)$.

We have therefore shown that:

$$\mathbb{E}[\hat{\mathbf{u}}^T \hat{\mathbf{u}} | \mathbf{X}] = (n - K) \sigma^2.$$

Therefore, it immediately follows that a conditionally unbiased estimator of σ^2 is:

$$\mathbb{E}\left[\frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{(n - K)} | \mathbf{X}\right] = \sigma^2$$

We can then use the law of iterated expectations to show that this is in fact unconditionally unbiased:

$$\mathbb{E}\left[\frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{(n - K)}\right] = \mathbb{E}\left[\mathbb{E}\left[\frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{(n - K)} | \mathbf{X}\right]\right] = \mathbb{E}[\sigma^2] = \sigma^2.$$

Therefore, the OLS estimator

$$\hat{\sigma}_{OLS}^2 = \frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{(n - K)}$$

is an unbiased estimator of the variance σ^2 . □

7.5 Interpretation of Coefficients in Regression Models

Suppose we had 2 regressors so the **multiple linear regression model**:

$$\mathbf{y} = \mathbf{X}_1 \beta_1 + \mathbf{X}_2 \beta_2 + \mathbf{u} \tag{128}$$

How can we interpret the coefficients β_1 and β_2 ? We can interpret it as the result of running two separate regressions which we shall outline below.

First, suppose we had the **simple linear regression**:

$$\mathbf{y} = \mathbf{X}_1 \gamma + r_y. \tag{129}$$

We then have that the OLS estimator $\hat{\gamma}$ would be:

$$\hat{\gamma} = (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{y}. \tag{130}$$

Now, also suppose we ran the regression:

$$\mathbf{X}_2 = \mathbf{X}_1 \delta + r_2. \tag{131}$$

We would have that the OLS estimator δ would be:

$$\hat{\delta}_{OLS} = (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_2. \quad (132)$$

Go back to Eq. (130) and plug in the multiple linear regression model:

$$\begin{aligned} \hat{\gamma}_{OLS} &= (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{y} \\ &= (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T [\mathbf{X}_1 \hat{\beta}_1 + \mathbf{X}_2 \hat{\beta}_2 + \hat{\mathbf{u}}] \\ &= (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_1 \hat{\beta}_1 + (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_2 \hat{\beta}_2 + (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \hat{\mathbf{u}} \\ &=_{(1)} \hat{\beta}_1 + \hat{\delta}_{OLS} \hat{\beta}_2 + (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \hat{\mathbf{u}} \\ &= \hat{\beta}_1 + \hat{\delta}_{OLS} \hat{\beta}_2 + 0 \\ &= \hat{\beta}_1 + \hat{\delta}_{OLS} \hat{\beta}_2 \end{aligned}$$

where $=_{(1)}$ comes from Eq. (132)

We therefore have that:

$$\hat{\gamma}_{OLS} = \hat{\beta}_1 + \hat{\delta}_{OLS} \hat{\beta}_2 \quad (133)$$

We can see that the OLS estimator in the multiple linear regression does not equal the OLS estimator in the simple linear regression $\hat{\beta}_1 \neq \hat{\gamma}_{OLS}$ unless either:

1. $\hat{\beta}_2 = 0$, so that $\mathbf{X}_2$ has no estimated partial relationship with $\mathbf{y}$;
2. $\hat{\delta}_{OLS} = 0$, so that $\mathbf{X}_1$ and $\mathbf{X}_2$ are orthogonal in the sample;

We can summarise it as this: the estimated coefficient in the multiple regression model considers the relationship between $\mathbf{y}$ and $\mathbf{X}_1$, controlling for the effect of $\mathbf{X}_2$.

7.6 Inference in Regression Models

We are now interested in making statistical statements about our OLS point estimates $\hat{\beta}$. However, in order to do so, we need further assumptions. These assumptions can take one of two forms:

1. Make the assumption that $\mathbf{u}|\mathbf{X}$ has a normal distribution and this will allow us to derive exact, **finite** sample distribution properties of $\hat{\beta}_{OLS}|\mathbf{X}$;
2. Make **no** assumption on $\mathbf{u}|\mathbf{X}$ but use **asymptotics** to derive distribution properties of $\hat{\beta}_{OLS}|\mathbf{X}$.

In this section, we will look at the properties under finite samples by making the assumption of normality. We will look at asymptotics in the next section.

7.6.1 Classical Linear Regression Model with Normally Distributed Error Terms

We state the assumption of normality on the error term.

Assumption 197 (Classical Linear Regression Model with Normally Distributed Error Terms)

The classical linear regression model with normally distributed error terms can be stated as:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u} \quad (134)$$

$$\mathbf{u}|\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I}) \quad (135)$$

$$\mathbf{X} \text{ has full rank with probability one.} \quad (136)$$

Alternatively, we can write:

$$\mathbf{y}|\mathbf{X} \sim \mathcal{N}(\mathbf{X}\boldsymbol{\beta}, \sigma^2 \mathbf{I}) \quad (137)$$

$$\mathbf{X} \text{ has full rank with probability one.} \quad (138)$$

We now get a useful result from this assumption of normality on the error term.

Theorem 198 (Conditional Normality of the OLS Estimator)

The conditional distribution of the OLS estimator is:

$$\hat{\boldsymbol{\beta}}_{OLS}|\mathbf{X} \sim \mathcal{N}(\boldsymbol{\beta}, \sigma^2(\mathbf{X}^T \mathbf{X})^{-1}) \quad (139)$$

Proof. We know we can write the OLS estimator as:

$$\hat{\boldsymbol{\beta}}_{OLS} = \boldsymbol{\beta} + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{u}$$

We know that $\mathbf{u}|\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$ by assumption.

We know that if $\mathbf{t} \sim \mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ and $\mathbf{z} = \mathbf{A} + \mathbf{B}\mathbf{t}$ for fixed matrices $\mathbf{A}, \mathbf{B}$, then $\mathbf{z} \sim \mathcal{N}(\mathbf{A} + \mathbf{B}\boldsymbol{\mu}, \mathbf{B}\boldsymbol{\Sigma}\mathbf{B}^T)$. Conditional on $\mathbf{X}$, the relevant matrices are fixed, so we can apply this to our OLS estimator expression:

$$\begin{aligned} \boldsymbol{\beta} + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{u} &\stackrel{D}{=} \boldsymbol{\beta} + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I}) \\ &= \mathcal{N}(\boldsymbol{\beta} + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \cdot \mathbf{0}, (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \sigma^2 \mathbf{I} \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1}) \\ &= \mathcal{N}(\boldsymbol{\beta}, \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1}) \\ &= \mathcal{N}(\boldsymbol{\beta}, \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1}) \end{aligned}$$

□

Looking at

$$\hat{\boldsymbol{\beta}}_{OLS}|\mathbf{X} \sim \mathcal{N}(\boldsymbol{\beta}, \sigma^2(\mathbf{X}^T \mathbf{X})^{-1})$$

we now want to construct test statistics.

Proposition 199 (Marginal Distribution of OLS Estimator)

Suppose that

$$\hat{\boldsymbol{\beta}}_{OLS}|\mathbf{X} \sim \mathcal{N}(\boldsymbol{\beta}, \sigma^2(\mathbf{X}^T \mathbf{X})^{-1})$$

then we have that:

$$\hat{\beta}_k|\mathbf{X} \sim \mathcal{N}(\beta_k, v_{kk}) \quad \text{for } k = 1, \dots, K \quad (140)$$

where $\hat{\beta}_k$ is the k^{th} element of $\hat{\boldsymbol{\beta}}_{OLS}$ and $v_{kk} = \text{Var}[\hat{\beta}_k|\mathbf{X}]$ is the k^{th} diagonal element of $\text{Var}[\hat{\boldsymbol{\beta}}_{OLS}|\mathbf{X}] = \sigma^2(\mathbf{X}^T \mathbf{X})^{-1}$.

Proof. Immediately follows by the marginal distribution property of multivariate normal distributions! $\square$

We have derived a conditional distribution result regarding the OLS estimator. Can we strengthen this to an unconditional result? No! However, we can if we standardise our OLS estimator. Therefore, we will now show why we standardise our OLS estimators when conducting hypothesis testing.

Theorem 200 (Unconditional Distribution of Standardised OLS Estimator)

The **unconditional distribution** of the standardised OLS estimator is given by:

$$z_k = \left(\frac{\hat{\beta}_k - \beta_k}{\sqrt{v_{kk}}} \right) \sim \mathcal{N}(0, 1) \quad (141)$$

Proof. First, note that

$$z_k = \left(\frac{\hat{\beta}_k - \beta_k}{\sqrt{v_{kk}}} \right) = \left(\frac{-\beta_k}{\sqrt{v_{kk}}} \right) + \left(\frac{1}{\sqrt{v_{kk}}} \right) \hat{\beta}_k$$

We know that $\hat{\beta}_k|\mathbf{X} \sim \mathcal{N}(\beta_k, v_{kk})$ so we then have

$$\begin{aligned} \left(\frac{-\beta_k}{\sqrt{v_{kk}}} \right) + \left(\frac{1}{\sqrt{v_{kk}}} \right) \hat{\beta}_k &\stackrel{D}{=} \left(\frac{-\beta_k}{\sqrt{v_{kk}}} \right) + \left(\frac{1}{\sqrt{v_{kk}}} \right) \mathcal{N}(\beta_k, v_{kk}) \\ &= \mathcal{N}\left(\frac{-\beta_k}{\sqrt{v_{kk}}} + \frac{\beta_k}{\sqrt{v_{kk}}}, \frac{v_{kk}}{v_{kk}} \right) \\ &= \mathcal{N}(0, 1). \end{aligned}$$

We therefore have that the conditional distribution is:

$$z_k|\mathbf{X} \sim \mathcal{N}(0, 1) \quad (142)$$

However, notice that the conditional distribution does not contain $\mathbf{X}$ at all so therefore, we can actually conclude that the **unconditional distribution** is:

$$z_k \sim \mathcal{N}(0, 1)$$

$\square$

Remark 201. Note that we can derive an **unconditional distribution** result

$$z_k \sim \mathcal{N}(0, 1)$$

since the conditional distribution did **not** depend on $\mathbf{X}$:

$$z_k|\mathbf{X} \sim \mathcal{N}(0, 1).$$

However, we cannot do the same for the conditional distribution of the OLS estimator:

$$\hat{\beta}_k | \mathbf{X} \sim \mathcal{N}(\beta_k, v_{kk}).$$

This is because the conditional distribution depends on $\mathbf{X}$ via the v_{kk} term.

In practice, we do not know what v_{kk} is. We therefore need to estimate it. We can use the unbiased OLS estimator:

$$\hat{\sigma}_{OLS}^2 = \frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{n - K}$$

to estimate σ^2 and then estimate $\text{Var}[\hat{\beta}_{OLS} | \mathbf{X}]$ using:

$$\widehat{\text{Var}}[\hat{\beta}_{OLS} | \mathbf{X}] = \hat{\sigma}_{OLS}^2 (\mathbf{X}^T \mathbf{X})^{-1}.$$

We denote $\hat{v}_{kk}$ to be the k^{th} element on the main diagonal of $\widehat{\text{Var}}[\hat{\beta}_{OLS} | \mathbf{X}]$.

Definition 202 (t-Statistic)

The **t-statistic** is given by:

$$t_k = \left(\frac{\hat{\beta}_k - \beta_k}{\sqrt{\hat{v}_{kk}}} \right) = \left(\frac{\hat{\beta}_k - \beta_k}{se_k} \right) \quad (143)$$

whereby $se_k = \sqrt{\hat{v}_{kk}}$ is called the **standard error** of the estimated parameter $\hat{\beta}_k$.

A big issue is now that we previously knew the distribution

$$z_k = \left(\frac{\hat{\beta}_k - \beta_k}{\sqrt{v_{kk}}} \right) \sim \mathcal{N}(0, 1)$$

but since we replaced $\sqrt{v_{kk}}$ with se_k , this changes the finite sample distribution of the statistic. We now need to find it.

Theorem 203 (Finite Sample Unconditional Distribution of t-statistic)

The (finite sample) unconditional distribution of the t-test statistic is given by:

$$t_k = \left(\frac{\hat{\beta}_k - \beta_k}{\sqrt{\hat{v}_{kk}}} \right) = \left(\frac{\hat{\beta}_k - \beta_k}{se_k} \right) \sim t(n - K) \quad (144)$$

where $t(n - K)$ is the t-distribution with $n - K$ degrees of freedom.

Proof. (The proof for this is quite long but very informative). First, we can rewrite

$$\begin{aligned} t_k &= \left(\frac{\hat{\beta}_k - \beta_k}{\sqrt{\hat{v}_{kk}}} \right) \\ &= \left(\frac{\sqrt{v_{kk}}}{\sqrt{\hat{v}_{kk}}} \right) \left(\frac{\hat{\beta}_k - \beta_k}{\sqrt{v_{kk}}} \right) \\ &= \left(\frac{\sqrt{v_{kk}}}{\sqrt{\hat{v}_{kk}}} \right) z_k \end{aligned}$$

where we have already shown that $z_k \sim \mathcal{N}(0, 1)$. We are interested in the expression

$$\left(\frac{\sqrt{v_{kk}}}{\sqrt{\hat{v}_{kk}}} \right).$$

First, recall that

$$\text{Var}[\hat{\boldsymbol{\beta}}_{OLS}|\mathbf{X}] = \sigma^2(\mathbf{X}^T\mathbf{X})^{-1}$$

If we let $\mathbf{Q} = (\mathbf{X}^T\mathbf{X})$ and q_{kk} to be the k -th diagonal element of $\mathbf{Q}^{-1}$, we have that

$$\sqrt{v_{kk}} = \sigma\sqrt{q_{kk}} \quad \text{and} \quad \sqrt{\hat{v}_{kk}} = \hat{\sigma}_{OLS}\sqrt{q_{kk}}$$

We therefore have that

$$\left(\frac{\sqrt{v_{kk}}}{\sqrt{\hat{v}_{kk}}}\right) = \frac{\sigma\sqrt{q_{kk}}}{\hat{\sigma}_{OLS}\sqrt{q_{kk}}} = \frac{\sigma}{\hat{\sigma}_{OLS}}. \quad (145)$$

We will keep this in mind later. We now look at the inverse as it is easier to derive the distribution of the inverse:

$$\frac{\hat{\sigma}_{OLS}^2}{\sigma^2} = \left(\frac{1}{\sigma^2}\right)\left(\frac{\hat{\mathbf{u}}^T\hat{\mathbf{u}}}{n-K}\right) = \left(\frac{1}{n-K}\right)\left(\frac{\hat{\mathbf{u}}^T\hat{\mathbf{u}}}{\sigma^2}\right) = \frac{w_0}{n-K} \quad (146)$$

where we defined $w_0 := \hat{\mathbf{u}}^T\hat{\mathbf{u}}/\sigma^2$. We make the following claim:

Claim 204.

$$w_0|\mathbf{X} \stackrel{D}{=} \left(\frac{\hat{\mathbf{u}}^T\hat{\mathbf{u}}}{\sigma^2}\right)|\mathbf{X} \sim \chi^2(n-K).$$

Proof. By assumption of normality, we had that:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u} \quad \text{where } \mathbf{u}|\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \sigma^2\mathbf{I})$$

Define $\mathbf{z}|\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ so we can rewrite this as:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \sigma\mathbf{z}$$

Now, recall that the fitted residual is the annihilator matrix $\mathbf{M}$ applied to the observed variable:

$$\begin{aligned} \hat{\mathbf{u}} &= \mathbf{M}\mathbf{y} = \mathbf{M}[\mathbf{X}\boldsymbol{\beta} + \sigma\mathbf{z}] \\ &\stackrel{(1)}{=} \mathbf{M}\mathbf{X}\boldsymbol{\beta} + \mathbf{M}\sigma\mathbf{z} \\ &= \mathbf{M}\sigma\mathbf{z} \end{aligned}$$

where $\stackrel{(1)}{=}$ uses the fact that $\mathbf{M}\mathbf{X} = \mathbf{0}$. We therefore have that the residual sum of squares:

$$\begin{aligned} \hat{\mathbf{u}}^T\hat{\mathbf{u}} &= (\mathbf{M}\sigma\mathbf{z})^T\mathbf{M}\sigma\mathbf{z} = \sigma\mathbf{z}^T\mathbf{M}^T\mathbf{M}\sigma\mathbf{z} \\ &\stackrel{(1)}{=} \sigma\mathbf{z}^T\mathbf{M}\sigma\mathbf{z} \\ &\stackrel{(2)}{=} \sigma^2\mathbf{z}^T\mathbf{M}\mathbf{z} \end{aligned}$$

whereby $\stackrel{(1)}{=}$ uses the fact that $\mathbf{M}$ is idempotent and $\stackrel{(2)}{=}$ uses the fact that σ is a scalar. We therefore have that:

$$w_0 = \left(\frac{\hat{\mathbf{u}}^T\hat{\mathbf{u}}}{\sigma^2}\right) = \left(\frac{\sigma^2\mathbf{z}^T\mathbf{M}\mathbf{z}}{\sigma^2}\right) = \mathbf{z}^T\mathbf{M}\mathbf{z}$$

whereby $\mathbf{z}|\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ and $\text{rank}(\mathbf{M}) = n - K$ by the properties of annihilator matrices. It is a known property of quadratic forms that since $\mathbf{z}|\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ and $\mathbf{M}$ is a symmetric and idempotent matrix, we have that the scalar:

$$\mathbf{z}^T\mathbf{M}\mathbf{z}|\mathbf{X} \sim \chi^2(\text{rank}(\mathbf{M})) = \chi^2(n-K).$$

We have shown therefore shown our intermediate claim that

$$w_0|\mathbf{X} \stackrel{D}{=} \left(\frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{\sigma^2} \right) | \mathbf{X} \sim \chi^2(n-K).$$

□

We now go back to our earlier expression for the t-test statistic

$$\begin{aligned} t_k &= \left(\frac{\sqrt{V_{kk}}}{\sqrt{\hat{V}_{kk}}} \right) z_k \\ &\stackrel{(1)}{=} \left(\frac{\sigma}{\hat{\sigma}_{OLS}} \right) z_k \\ &= \frac{z_k}{\hat{\sigma}_{OLS}/\sigma} \\ &= \frac{z_k}{\sqrt{\hat{\sigma}_{OLS}^2/\sigma^2}} \\ &\stackrel{(2)}{=} \frac{z_k}{\sqrt{w_0/(n-K)}} \end{aligned}$$

whereby $\stackrel{(1)}{=}$ uses Eq. (145) and $\stackrel{(2)}{=}$ uses Eq. (146). To recap, we now have the expression for our test statistic as:

$$t_k = \frac{z_k}{\sqrt{w_0/(n-K)}} \quad (147)$$

whereby we know that

$$z_k|\mathbf{X} \sim \mathcal{N}(0,1) \quad \text{and} \quad w_0|\mathbf{X} \sim \chi^2(n-K).$$

We have a normal distribution divided by a χ^2 -distribution. To show that this is a t-distribution, we need to show that the numerator and the denominator are independent. We shall show this in our next claim.

Claim 205. z_k and w_0 are independent after conditioning on $\mathbf{X}$.

Proof. Let us take a look at our two expressions again:

$$z_k = \frac{\hat{\beta}_k - \beta_k}{\sqrt{V_{kk}}} \quad \text{and} \quad w_0 = \left(\frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{\sigma^2} \right)$$

We see that z_k is a function of $\hat{\beta}_{OLS}$ and w_0 is a function of $\hat{\mathbf{u}}$. We will instead show that any function of $\hat{\beta}_{OLS}$ and any function of $\hat{\mathbf{u}}$ are independent, which will imply the claim we are making.

First, write the OLS expression:

$$\hat{\beta}_{OLS} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T [\mathbf{X}\boldsymbol{\beta} + \sigma \mathbf{z}] = \boldsymbol{\beta} + \sigma (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{z}$$

where $\mathbf{z}|\mathbf{X} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$. We can rewrite this as:

$$\hat{\beta}_{OLS} - \boldsymbol{\beta} = \sigma (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{z}.$$

Now, consider the following random vector (these may seem like they come out of nowhere but will prove to be useful):

$$\begin{aligned} d_1 &= \frac{\hat{\mathbf{u}}}{\sigma} = \frac{\sigma \mathbf{M} \mathbf{z}}{\sigma} = \mathbf{M} \mathbf{z} \\ d_2 &= \frac{\hat{\beta}_{OLS} - \boldsymbol{\beta}}{\sigma} = \frac{\sigma (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{z}}{\sigma} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{z} \end{aligned}$$

Both $\mathbf{M}$ and $(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$ are known when conditioned on $\mathbf{X}$. Therefore, we have that:

$$d_1 | \mathbf{X} \sim \mathcal{N}(\cdot, \cdot) \quad d_2 | \mathbf{X} \sim \mathcal{N}(\cdot, \cdot)$$

as $\mathbf{z} | \mathbf{X} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$. Furthermore, we have that $\mathbf{d} | \mathbf{X}$ is normally distributed whereby $\mathbf{d} = (d_1, d_2)^T$. We can then show that d_1 and d_2 are uncorrelated by first showing that:

$$\begin{aligned} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{M} &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T [\mathbf{I} - \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T] \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T - (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T - (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T = \mathbf{0}. \end{aligned}$$

It follows that

$$\text{Cov}[d_2, d_1 | \mathbf{X}] = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{M} \text{Var}[\mathbf{z} | \mathbf{X}] = 0.$$

This means that d_1, d_2 are uncorrelated conditioned on $\mathbf{X}$. This then implies that d_1, d_2 are independent since this is a multivariate normal distribution. We can then rewrite:

$$\begin{aligned} d_1 &= \frac{\hat{\mathbf{u}}}{\sigma} \leftrightarrow \hat{\mathbf{u}} = \sigma d_1 \\ d_2 &= \frac{\hat{\boldsymbol{\beta}}_{OLS} - \boldsymbol{\beta}}{\sigma} \leftrightarrow \hat{\boldsymbol{\beta}}_{OLS} = \boldsymbol{\beta} + \sigma d_2 \end{aligned}$$

whereby $\boldsymbol{\beta}, \sigma$ are non-stochastic. This implies that conditioned on $\mathbf{X}$, $\hat{\boldsymbol{\beta}}_{OLS}$ and $\hat{\mathbf{u}}$ are therefore independent. Hence, we can conclude that conditional on $\mathbf{X}$, any function of $\hat{\mathbf{u}}$ and any function of $\hat{\boldsymbol{\beta}}_{OLS}$ are independent. $\square$

Therefore,

$$t_k = \left(\frac{\hat{\beta}_k - \beta_k}{\sqrt{\hat{v}_{kk}}} \right) = \frac{z_k}{\sqrt{w_0/(n-K)}}$$

is the ratio of a standard normal random variable and the square root of an independent χ^2 random variable divided by its degrees of freedom. By Proposition 55, we can therefore conclude that the conditional distribution of the t-statistic is:

$$t_k | \mathbf{X} \sim t(n-K) \quad (148)$$

As the conditional distribution of t_k does **not** depend on $\mathbf{X}$, this means that this holds for all $\mathbf{X}$ and we can in fact state that the **unconditional distribution** is:

$$t_k \sim t(n-K). \quad (149)$$

$\square$

Proposition 206 (Asymptotic Normality of t-statistic)

Denote $\nu = n - K$ and let the finite-sample t-test statistic satisfy $t_k \sim t(\nu)$. Then

$$t_k \xrightarrow{D} \mathcal{N}(0, 1) \quad (150)$$

as $n \rightarrow \infty$.

Proof. Write the t-statistic as

$$t_k = \frac{z}{\sqrt{w/\nu}}$$

whereby $z \sim \mathcal{N}(0, 1)$ and $w \sim \chi^2(\nu)$.

Looking at the denominator, we know that a χ^2 random variable can be thought of as:

$$w = \sum_{i=1}^{\nu} z_i^2$$

where $z_i \sim \mathcal{N}(0, 1)$ for $i = 1, 2, \dots, \nu$ are independent standard normal random variables. First, we note that

$$\mathbb{E}[z_i] = 0$$

and

$$\begin{aligned} \text{Var}[z_i] &= \mathbb{E}[z_i^2] - \mathbb{E}[z_i]^2 \leftrightarrow 1 = \mathbb{E}[z_i^2] - 0^2 \\ &\Rightarrow \mathbb{E}[z_i^2] = 1 \end{aligned}$$

Therefore, computing the first moment of w , we have:

$$\begin{aligned} \mathbb{E}[w] &= \mathbb{E}\left[\sum_{i=1}^{\nu} z_i^2\right] \\ &= \nu \mathbb{E}[z_i^2] \\ &=_{(1)} \nu \times 1 \\ &= \nu. \end{aligned}$$

where $=_{(1)}$ uses the fact that $\mathbb{E}[z_i^2] = 1$.

Then, we have that in the denominator:

$$\mathbb{E}\left[\frac{w}{\nu}\right] = \frac{1}{\nu} \mathbb{E}[w] = 1.$$

We have verified that z_i are independent and has finite first and second moments. We can therefore apply the **weak law of large numbers** to the denominator:

$$\frac{w}{\nu} = \frac{1}{\nu} w = \frac{1}{\nu} \sum_{i=1}^{\nu} z_i^2 \xrightarrow{P} \mathbb{E}[z_i^2] = 1.$$

By the **continuous mapping theorem**, we therefore have that:

$$\sqrt{\frac{w}{\nu}} \xrightarrow{P} \sqrt{1} = 1.$$

Going back to our t-statistic t_k , we can therefore apply the **Slutsky's theorem** and get our desired result:

$$t_k = \frac{z}{\sqrt{w/\nu}} \xrightarrow{D} z \sim \mathcal{N}(0, 1).$$

This is the desired convergence-in-distribution result.

□

7.6.2 Hypothesis Testing

We have shown that:

$$t_k = \left(\frac{\hat{\beta}_k - \beta_k}{\sqrt{\hat{v}_{kk}}} \right) = \left(\frac{\hat{\beta}_k - \beta_k}{se_k} \right) \sim t(n - K).$$

Suppose we are interested in the hypothesis test that

$$H_0 : \beta_k = 0.$$

Then, the **t-ratio** is

$$t_k = \left(\frac{\hat{\beta}_k - 0}{se_k} \right) = \left(\frac{\hat{\beta}_k}{se_k} \right) \sim t(n - K)$$

under $H_0 : \beta_k = 0$, which is simply the ratio of the OLS estimate of the coefficient of $\mathbf{x}_k$ to its standard error.

We now discuss hypothesis testing when talking about linear combination of parameters. Recall that we have K parameters in $\boldsymbol{\beta}$. Suppose we want to test $p \geq 1$ linear restrictions. We would then need a $p \times K$ matrix $\mathbf{H}$ whereby each row of $\mathbf{H}$ contains the coefficients of one of the linear restrictions. We can then store these linear restrictions in a $p \times 1$ vector $\boldsymbol{\theta}$. We therefore have that:

$$\boldsymbol{\theta} = \mathbf{H}\boldsymbol{\beta} \quad (151)$$

Example 207

Suppose we have a Cobb-Douglas production function of the form:

$$y_i = \beta_1 + \beta_2 x_{2i} + \beta_3 x_{3i} + u_i$$

where y_i is the log of value added for observation i , x_{2i} is the log of labour input for observation i and x_{3i} is the log of capital input for observations i . We may be interested in testing the constant returns to scale restriction:

$$H_0 : \beta_2 + \beta_3 = 1.$$

This is only one restriction. Therefore, we need a 1×3 matrix $\mathbf{H}$.

We would have that the $\mathbf{H}$ matrix would then be:

$$\mathbf{H} = \begin{bmatrix} 0 & 1 & 1 \end{bmatrix}$$

which will give us

$$\boldsymbol{\theta} = \mathbf{H}\boldsymbol{\beta} = \begin{bmatrix} \beta_2 + \beta_3 \end{bmatrix}.$$

We now derive a distribution result for a general test statistic.

Theorem 208 (Distribution of F-statistic)

Consider $p \geq 1$ linear restrictions. We have the $p \times 1$ vector $\boldsymbol{\theta} = \mathbf{H}\boldsymbol{\beta}$ for some non-stochastic $p \times K$ matrix $\mathbf{H}$ with $\text{rank}(\mathbf{H}) = p$ (full row rank). Estimate $\boldsymbol{\theta}$ by using the OLS $\hat{\boldsymbol{\theta}}_{OLS} = \mathbf{H}\hat{\boldsymbol{\beta}}_{OLS}$. The (scalar) F-test statistic has the form:

$$v = \left(\frac{1}{p\hat{\sigma}_{OLS}^2} \right) (\hat{\boldsymbol{\theta}}_{OLS} - \boldsymbol{\theta})^T D (\hat{\boldsymbol{\theta}}_{OLS} - \boldsymbol{\theta}) \quad (152)$$

$$= \left(\frac{1}{p} \right) (\hat{\boldsymbol{\theta}}_{OLS} - \boldsymbol{\theta})^T \left[\widehat{\text{Var}}[\hat{\boldsymbol{\theta}}_{OLS} | \mathbf{X}] \right]^{-1} (\hat{\boldsymbol{\theta}}_{OLS} - \boldsymbol{\theta}) \sim F(p, n - K) \quad (153)$$

whereby $D^{-1} = \mathbf{H}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{H}^T$.

Proof. (This is another lengthy proof but also very informative.)

We have $\hat{\theta}_{OLS} = H\hat{\beta}_{OLS}$. Under the assumptions of the classical linear regression model, we have that:

$$\mathbb{E}[\hat{\theta}_{OLS}|\mathbf{X}] = H\mathbb{E}[\hat{\beta}_{OLS}|\mathbf{X}] = H\boldsymbol{\beta} = \boldsymbol{\theta}$$

which shows that $\hat{\theta}_{OLS}$ is an unbiased estimator of $\boldsymbol{\theta}$.

For the variance, we have that

$$\begin{aligned} \text{Var}[\hat{\theta}_{OLS}|\mathbf{X}] &= H\text{Var}[\hat{\beta}_{OLS}|\mathbf{X}]H^T \\ &= H[\sigma^2(\mathbf{X}^T\mathbf{X})^{-1}]H^T \\ &= \sigma^2 H(\mathbf{X}^T\mathbf{X})^{-1}H^T \\ &= \sigma^2 D^{-1} \end{aligned}$$

whereby we define $D^{-1} = H(\mathbf{X}^T\mathbf{X})^{-1}H^T$. However, we don't know what σ^2 is so we can estimate it with $\hat{\sigma}_{OLS}^2$ to get:

$$\widehat{\text{Var}}[\hat{\theta}_{OLS}|\mathbf{X}] = \hat{\sigma}_{OLS}^2 D^{-1}. \quad (154)$$

With normally distributed error terms, we can also conclude that:

$$\hat{\theta}_{OLS}|\mathbf{X} \sim \mathcal{N}(\boldsymbol{\theta}, \sigma^2 D^{-1}) \quad (155)$$

We make the first claim.

Claim 209. *The scalar random variable*

$$w|\mathbf{X} = (\hat{\theta}_{OLS} - \boldsymbol{\theta})^T [\sigma^2 D^{-1}]^{-1} (\hat{\theta}_{OLS} - \boldsymbol{\theta})|\mathbf{X} \sim \chi^2(p)$$

Proof. We know our result about quadratic forms of normal random vectors that if we have a random variable $\mathbf{X}$ which has a multivariate normal distribution $\mathbf{X} \sim \mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ whereby $\mathbf{X}$ has dimension d and the covariance matrix $\boldsymbol{\Sigma}$ has full rank d then, we have that the following quadratic form of $\mathbf{X}$ has the distribution:

$$[\mathbf{X} - \boldsymbol{\mu}]^T \boldsymbol{\Sigma}^{-1} [\mathbf{X} - \boldsymbol{\mu}] \sim \chi^2(d)$$

which is a χ^2 distribution with d degrees of freedom.

We already know that $\hat{\theta}_{OLS}|\mathbf{X} \sim \mathcal{N}(\boldsymbol{\theta}, \sigma^2 D^{-1})$ and hence, we can apply our result immediately to the scalar:

$$w = (\hat{\theta}_{OLS} - \boldsymbol{\theta})^T [\sigma^2 D^{-1}]^{-1} (\hat{\theta}_{OLS} - \boldsymbol{\theta})$$

□

We have previously shown that w is a conditional χ^2 random variable:

$$w|\mathbf{X} = \left(\frac{1}{\sigma^2}\right) (\hat{\theta}_{OLS} - \boldsymbol{\theta})^T D (\hat{\theta}_{OLS} - \boldsymbol{\theta})|\mathbf{X}$$

We can estimate the variance to get:

$$\hat{w} = \left(\frac{1}{\hat{\sigma}_{OLS}^2}\right) (\hat{\theta}_{OLS} - \boldsymbol{\theta})^T D (\hat{\theta}_{OLS} - \boldsymbol{\theta})$$

We then redefine to get the variable:

$$v = \frac{\hat{w}}{p} = \left(\frac{1}{p\hat{\sigma}_{OLS}^2}\right) (\hat{\theta}_{OLS} - \boldsymbol{\theta})^T D (\hat{\theta}_{OLS} - \boldsymbol{\theta})$$

We then write:

$$\begin{aligned} v &= \frac{\hat{w}}{p} \\ &=_{(1)} \left(\frac{\sigma^2}{\hat{\sigma}_{OLS}^2} \right) \left(\frac{w}{p} \right) \\ &= \left(\frac{1}{p \hat{\sigma}_{OLS}^2} \right) (\hat{\theta}_{OLS} - \theta)^T D (\hat{\theta}_{OLS} - \theta) \end{aligned}$$

where $=_{(1)}$ arises from re-expressing w as $\hat{w}$. We have shown before that

$$\frac{\hat{\sigma}_{OLS}^2}{\sigma^2} = \frac{w_0}{n - K}$$

whereby $w_0 = \frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{\sigma^2}$ and that

$$w_0 | \mathbf{X} \sim \chi^2(n - K).$$

Therefore,

$$\begin{aligned} v &= \left(\frac{\sigma^2}{\hat{\sigma}_{OLS}^2} \right) \left(\frac{w}{p} \right) \\ &= \frac{(w/p)}{(\hat{\sigma}_{OLS}^2/\sigma^2)} \\ &= \frac{(w/p)}{(w_0/(n - K))} \end{aligned}$$

is the ratio of two chi-squared random variables, each divided by their respective degrees of freedom $w | \mathbf{X} \sim \chi^2(p)$ and $w_0 | \mathbf{X} \sim \chi^2(n - K)$.

Moreover, we have that

$$w = \left(\frac{1}{\sigma^2} \right) (\hat{\theta}_{OLS} - \theta)^T D (\hat{\theta}_{OLS} - \theta)$$

is a function of $\hat{\beta}_{OLS}$ whilst

$$w_0 = \frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{\sigma^2}$$

is a function of $\hat{\mathbf{u}}$. We have previously shown that conditioned on $\mathbf{X}$, any function of $\hat{\mathbf{u}}$ and any function of $\hat{\beta}_{OLS}$ are independent.

Therefore,

$$v = \frac{(w/p)}{(w_0/(n - K))}$$

is the ratio of two independent chi-squared random variables, each divided by their respective degrees of freedom.

We can conclude that the conditional distribution is:

$$v | \mathbf{X} \sim F(p, n - K) \quad (156)$$

Since this conditional distribution of $v | \mathbf{X}$ is the same distribution (F distribution with p and $n - K$ degrees of freedom) for any realisation of $\mathbf{X}$, we can therefore conclude that the unconditional distribution is the same:

$$v = \frac{(w/p)}{(w_0/(n - K))} \sim F(p, n - K) \quad (157)$$

□

If we want to test a hypothesis about p linear combinations of the K β parameters, we can formulate the null

hypothesis:

$$H_0 : H\boldsymbol{\beta} = \boldsymbol{\theta}^0$$

where H is a non-stochastic $p \times K$ matrix with full rank $\text{rank}(H) = p$.

The alternative hypothesis is therefore:

$$H_1 : H\boldsymbol{\beta} \neq \boldsymbol{\theta}^0$$

We use the estimator $\hat{\boldsymbol{\theta}}_{OLS} = H\hat{\boldsymbol{\beta}}_{OLS}$ to estimate our parameter of interest $\boldsymbol{\theta}$. We can then construct the scalar test statistic:

$$v = \left(\frac{1}{p}\right)(\hat{\boldsymbol{\theta}}_{OLS} - \boldsymbol{\theta}^0)^T \left[\widehat{\text{Var}}[\hat{\boldsymbol{\theta}}_{OLS}|\mathbf{X}]\right]^{-1} (\hat{\boldsymbol{\theta}}_{OLS} - \boldsymbol{\theta}^0)$$

whereby we have shown that in the classical linear regression model with normally distributed errors:

$$\widehat{\text{Var}}[\hat{\boldsymbol{\theta}}_{OLS}|\mathbf{X}] = \widehat{\text{Var}}[H\hat{\boldsymbol{\beta}}_{OLS}|\mathbf{X}] = H\widehat{\text{Var}}[\hat{\boldsymbol{\beta}}_{OLS}|\mathbf{X}]H^T = H\hat{\sigma}_{OLS}^2(\mathbf{X}^T\mathbf{X})^{-1}H^T.$$

If the null hypothesis is true, we have shown that the test statistic has a F-distribution:

$$v \sim F(p, n - K)$$

We ask whether the value of the test statistic v is an unlikely draw from this null distribution. We reject $H_0 : H\boldsymbol{\beta} = \boldsymbol{\theta}^0$ if the probability of drawing a test statistic with a value as high as v from an F-distribution with $(p, n - K)$ degrees of freedom is below our chosen level of significance, conventionally $\alpha = 5\%$. As the F-distribution has support only over the positive real numbers, we need to do a 1-sided upper tail test and see if the test statistic v lies above the upper 95% percentile of the $F(p, n - K)$ distribution.

Example 210

Suppose that $K = 3$ so we have that:

$$y_i = \beta_1 + \beta_2 x_{2i} + \beta_3 x_{3i} + u_i$$

If we want to test that $\beta_2 = 0$ and $\beta_3 = 0$, then our null hypothesis would be:

$$H_0 : H\boldsymbol{\beta} = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} \beta_1 \\ \beta_2 \\ \beta_3 \end{bmatrix} = \begin{bmatrix} \beta_2 \\ \beta_3 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix}$$

Using our OLS estimates, we have that:

$$\hat{\boldsymbol{\theta}}_{OLS} = H\hat{\boldsymbol{\beta}}_{OLS} = \begin{bmatrix} \hat{\beta}_2 \\ \hat{\beta}_3 \end{bmatrix}$$

The variance is given by:

$$\begin{aligned} \widehat{Var}[\hat{\boldsymbol{\theta}}_{OLS}|\mathbf{X}] &= H\widehat{Var}[\hat{\boldsymbol{\beta}}_{OLS}|\mathbf{X}]H^T \\ &= \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} \hat{v}_{11} & \hat{v}_{12} & \hat{v}_{13} \\ \hat{v}_{21} & \hat{v}_{22} & \hat{v}_{23} \\ \hat{v}_{31} & \hat{v}_{32} & \hat{v}_{33} \end{bmatrix} \begin{bmatrix} 0 & 0 \\ 1 & 0 \\ 0 & 1 \end{bmatrix} \\ &= \begin{bmatrix} \hat{v}_{22} & \hat{v}_{23} \\ \hat{v}_{23} & \hat{v}_{33} \end{bmatrix} \end{aligned}$$

We can then construct the test statistic to test $H_0 : \boldsymbol{\theta}^0 = \mathbf{0}$:

$$\begin{aligned} v &= \left(\frac{1}{p}\right) (\hat{\boldsymbol{\theta}}_{OLS} - \boldsymbol{\theta}^0)^T \left[\widehat{Var}[\hat{\boldsymbol{\theta}}_{OLS}|\mathbf{X}]\right]^{-1} (\hat{\boldsymbol{\theta}}_{OLS} - \boldsymbol{\theta}^0) \\ &= \left(\frac{1}{2}\right) \begin{bmatrix} \hat{\beta}_2 - 0 \\ \hat{\beta}_3 - 0 \end{bmatrix}^T \begin{bmatrix} \hat{v}_{22} & \hat{v}_{23} \\ \hat{v}_{23} & \hat{v}_{33} \end{bmatrix}^{-1} \begin{bmatrix} \hat{\beta}_2 - 0 \\ \hat{\beta}_3 - 0 \end{bmatrix} \\ &= \left(\frac{1}{2}\right) \begin{bmatrix} \hat{\beta}_2 & \hat{\beta}_3 \end{bmatrix} \begin{bmatrix} \hat{v}_{22} & \hat{v}_{23} \\ \hat{v}_{23} & \hat{v}_{33} \end{bmatrix}^{-1} \begin{bmatrix} \hat{\beta}_2 \\ \hat{\beta}_3 \end{bmatrix} \sim F(2, n - K) \end{aligned}$$

We therefore reject $H_0 : \boldsymbol{\theta} = \mathbf{0}$ at the 5% level of significance if the probability of drawing this value of v from the $F(2, n - K)$ distribution is less than 0.05 i.e. if v lies above the upper 95% percentile of the $F(2, n - K)$ distribution. For example, if $n - K = 30$, the upper 95% percentile of the $F(2, 30)$ distribution is 3.32 and values of the test statistic $v > 3.32$ would reject the null hypothesis at the 5% level.

7.6.3 Using Linear Transformations to Simplify Hypothesis Tests

We can reparameterise models so that we can do t-tests instead of F-tests. For example, suppose we had the model:

$$y_i = \beta_1 + \beta_2 x_{2i} + \beta_3 x_{3i} + u_i$$

and we wish to do a two-sided test on the null:

$$H_0 : \beta_2 + \beta_3 = 1.$$

We can subtract x_{3i} from both sides to get:

$$y_i - x_{3i} = \beta_1 + \beta_2 x_{2i} + (\beta_3 - 1)x_{3i} + u_i$$

We can then add and subtract $\beta_2 x_{3i}$ on the right hand side to get:

$$y_i - x_{3i} = \beta_1 + \beta_2(x_{2i} - x_{3i}) + (\beta_2 + \beta_3 - 1)x_{3i} + u_i$$

We can therefore run the regression of $(y_i - x_{3i})$ on the intercept, $(x_{2i} - x_{3i})$ and x_{3i} to get:

$$y_i - x_{3i} = \hat{\beta}_1 + \hat{\beta}_2(x_{2i} - x_{3i}) + \hat{\delta}x_{3i} + u_i$$

whereby $\hat{\delta}$ is an estimator of $(\beta_2 + \beta_3 - 1)$. Therefore, we can see that the simple t-test on:

$$H_0 : \hat{\delta} = 0 \quad \leftrightarrow \quad H_0 : \beta_2 + \beta_3 = 1$$

are equivalent!

Therefore, the general principle of this is to add and subtract terms so that one of your coefficients on your variables is the your original null hypothesis of interest.

7.6.4 Goodness of Fit

Recall that we can always decompose the observed variable y as the fitted value plus the fitted residual.

$$y = \hat{y} + \hat{u}$$

Proposition 211 (Fitted Values are Orthogonal to Fitted Residuals)

The fitted values are orthogonal to the fitted residuals from OLS:

$$\hat{y}^T \hat{u} = 0. \quad (158)$$

Proof.

$$\begin{aligned} \hat{y}^T \hat{u} &= (X\hat{\beta}_{OLS})^T \hat{u} \\ &= \hat{\beta}_{OLS}^T X^T \hat{u} \\ &=_{(1)} \hat{\beta}_{OLS}^T \cdot 0 = 0 \end{aligned}$$

whereby $=_{(1)}$ comes from the first order conditions used to obtain $\hat{\beta}_{OLS}$ whereby $X^T \hat{u} = 0$. □

Theorem 212 (Decomposition of Total Sum of Squares)

The total sum of squares (TSS) can be decomposed as the sum of the explained sum of squares (ESS) and the residual sum of squares (RSS) **provided** that the average of the sample residuals $\bar{\hat{u}}$ are 0:

$$\underbrace{\sum_{i=1}^n (y_i - \bar{y})^2}_{\text{TSS}} = \underbrace{\sum_{i=1}^n (\hat{y}_i - \bar{y})^2}_{\text{ESS}} + \underbrace{\sum_{i=1}^n \hat{u}_i^2}_{\text{RSS}} \quad (159)$$

Proof. First, we can write:

$$\begin{aligned} y^T y &= (\hat{y} + \hat{u})^T (\hat{y} + \hat{u}) \\ &= \hat{y}^T \hat{y} + \hat{y}^T \hat{u} + \hat{u}^T \hat{y} + \hat{u}^T \hat{u} \\ &= \hat{y}^T \hat{y} + 0 + 0 + \hat{u}^T \hat{u} \\ &= \hat{y}^T \hat{y} + \hat{u}^T \hat{u} \end{aligned}$$

which is equivalent to:

$$\sum_{i=1}^n y_i^2 = \sum_{i=1}^n \hat{y}_i^2 + \sum_{i=1}^n \hat{u}_i^2$$

We can subtract $n\bar{y}^2$ from both sides:

$$\begin{aligned} \sum_{i=1}^n y_i^2 - n\bar{y}^2 &= \sum_{i=1}^n \hat{y}_i^2 - n\bar{y}^2 + \sum_{i=1}^n \hat{u}_i^2 \\ \sum_{i=1}^n (y_i - \bar{y})^2 &= \underbrace{\sum_{i=1}^n \hat{y}_i^2 - n\bar{y}^2}_{(*)} + \sum_{i=1}^n \hat{u}_i^2 \end{aligned}$$

We are now interested in the (*) term. First, since we have that $y_i = \hat{y}_i + \hat{u}_i$, this implies:

$$\sum_{i=1}^n y_i = \sum_{i=1}^n \hat{y}_i + \sum_{i=1}^n \hat{u}_i$$

and dividing through by n, we get:

$$\bar{y} = \bar{\hat{y}} + \bar{\hat{u}}$$

By assumption, we said that the average of the fitted residuals is 0:

$$\bar{\hat{u}} = 0$$

which means:

$$\bar{y} = \bar{\hat{y}}.$$

We can then plug this into (*) to get:

$$\sum_{i=1}^n \hat{y}_i^2 - n\bar{y}^2 = \sum_{i=1}^n (\hat{y}_i - \bar{y})^2.$$

Therefore, we have arrived at our desired expression:

$$\underbrace{\sum_{i=1}^n (y_i - \bar{y})^2}_{\text{TSS}} = \underbrace{\sum_{i=1}^n (\hat{y}_i - \bar{y})^2}_{\text{ESS}} + \underbrace{\sum_{i=1}^n \hat{u}_i^2}_{\text{RSS}}$$

□

Remark 213. We made the seemingly random assumption that the average of the sample residuals is 0: $\bar{\hat{u}} = 0$. This in fact always hold if **the model includes an intercept term**.

We can take our expression:

$$\sum_{i=1}^n (y_i - \bar{y})^2 = \sum_{i=1}^n (\hat{y}_i - \bar{y})^2 + \sum_{i=1}^n \hat{u}_i^2$$

and divide by $\sum_{i=1}^n (y_i - \bar{y})^2$ on both sides to get:

$$1 = \frac{\sum_{i=1}^n (\hat{y}_i - \bar{y})^2}{\sum_{i=1}^n (y_i - \bar{y})^2} + \frac{\sum_{i=1}^n \hat{u}_i^2}{\sum_{i=1}^n (y_i - \bar{y})^2}$$

We can then re-arrange this to get:

$$\frac{\sum_{i=1}^n (\hat{y}_i - \bar{y})^2}{\sum_{i=1}^n (y_i - \bar{y})^2} = 1 - \frac{\sum_{i=1}^n \hat{u}_i^2}{\sum_{i=1}^n (y_i - \bar{y})^2}$$

Definition 214 (R-squared)

R-squared is given by:

$$R^2 = \frac{\sum_{i=1}^n (\hat{y}_i - \bar{y})^2}{\sum_{i=1}^n (y_i - \bar{y})^2} = 1 - \frac{\sum_{i=1}^n \hat{u}_i^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (160)$$

or equivalently:

$$R^2 = \frac{\text{ESS}}{\text{TSS}} = 1 - \frac{\text{RSS}}{\text{TSS}} \quad (161)$$

whereby $0 \leq R^2 \leq 1$ provided that $\bar{\hat{u}} = 0$.

R^2 measures the proportion of the total variation in the dependent variable that is accounted for by the model. In fact, R^2 measures the goodness of fit of a model with additional explanatory variables **relative** to the fit of a model that includes only an intercept term. Hence, we should not use this centred R^2 as a measure of fit for models that do not include an intercept term. With an intercept, R^2 is invariant to affine rescaling of the dependent variable and to reparameterisations of the regressors that preserve the column space of X . However, R^2 should **not** be used to compare the fit of models that have **different dependent variables**, such as levels and logarithms.

7.6.5 Using R^2 to test exclusion restrictions

An alternative way to test p exclusion restrictions is to compute the test statistic:

$$v = \left(\frac{n - K}{p} \right) \left(\frac{R_U^2 - R_R^2}{1 - R_U^2} \right) \sim F(p, n - K) \quad (162)$$

whereby R_U^2 is the R^2 from the unrestricted model and R_R^2 is the R^2 from the restricted model. Here, the null hypothesis is rejected if the exclusion of these p explanatory variables results in a sufficiently large fall in the R^2 goodness of fit measure, or in a sufficiently large deterioration in the fit of the model.

This test can be used to test whether do **all** $K - 1$ slope coefficients in the model are equal to 0. Here, the restricted model $y_i = \beta_1 + u_i$ has $R_R^2 = 0$. Therefore, the test statistic simplifies to:

$$v = \left(\frac{n - K}{K - 1} \right) \left(\frac{R^2}{1 - R^2} \right) \sim F(K - 1, n - K) \quad (163)$$

where $R^2 = R_U^2$ denotes the R^2 in the unrestricted model. This is the **F-test** for the model.

7.7 Asymptotic Properties in Regression Models

So far we have looked at regression in a finite sample setting. Under such a setting, we had to impose strong assumptions of normality on error terms. We can in fact relax this assumption of normality by instead considering asymptotics. We now want to look at properties of the OLS estimators that hold in large samples. That is, fixing the number of explanatory variables K fixed, we let the sample size n go to infinity. This approach of analysing properties actually hold under **weaker assumptions** than under finite sample results.

7.7.1 Consistency of OLS Estimator

We are interested in showing that the OLS estimator is consistent. We state the assumptions required for this.

Assumption 215 (Assumptions for Consistency)

The assumptions needed for consistency of the OLS estimator are:

1. $\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}$;
2. The data on $(y_i, \mathbf{x}_i^T)$ are independent and identically distributed for $i = 1, \dots, n$;
3. $\mathbb{E}[\mathbf{x}_i u_i] = 0$ for all $i = 1, \dots, n$;
4. The $K \times K$ matrix $M_{XX} = \mathbb{E}[\mathbf{x}_i \mathbf{x}_i^T]$ exists and is non-singular.

Theorem 216 (Consistency of OLS Estimator)

Under the assumptions of consistency, the OLS estimator is **consistent**

$$p \lim_{n \rightarrow \infty} \hat{\boldsymbol{\beta}}_{OLS} = \boldsymbol{\beta} \quad (164)$$

or equivalently:

$$\hat{\boldsymbol{\beta}}_{OLS} \xrightarrow{P} \boldsymbol{\beta}. \quad (165)$$

Proof. First, write the OLS estimator as:

$$\begin{aligned} \hat{\boldsymbol{\beta}}_{OLS} &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T (\mathbf{X}\boldsymbol{\beta} + \mathbf{u}) \\ &= (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X}\boldsymbol{\beta} + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{u} \\ &= \boldsymbol{\beta} + (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{u} \end{aligned}$$

We then multiply and divide by n :

$$\begin{aligned} \hat{\boldsymbol{\beta}}_{OLS} &= \boldsymbol{\beta} + n(\mathbf{X}^T \mathbf{X})^{-1} \left(\frac{1}{n} \right) \mathbf{X}^T \mathbf{u} \\ &= \boldsymbol{\beta} + \underbrace{\left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1}}_{(*)} \underbrace{\left(\frac{\mathbf{X}^T \mathbf{u}}{n} \right)}_{(**)} \end{aligned}$$

We first look at (*). By assumption (2), the observations $\mathbf{x}_i$ are i.i.d. Furthermore, by assumption (4), the second moment $\mathbb{E}[\mathbf{x}_i \mathbf{x}_i^T]$ exists. Therefore, by the **weak law of large numbers**, we have that:

$$\frac{\mathbf{X}^T \mathbf{X}}{n} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^T \xrightarrow{P} \mathbb{E}[\mathbf{x}_i \mathbf{x}_i^T] =_{(1)} M_{XX}$$

whereby $=_{(1)}$ comes from the definition of M_{XX} by definition. We can then apply the **continuous mapping theorem** to conclude that:

$$p \lim_{n \rightarrow \infty} \left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1} = (M_{XX})^{-1} =_{(1)} M_{XX}^{-1}$$

whereby in $=_{(1)}$, by assumption (4), we assumed that M_{XX} is non-singular so its inverse exists.

We now look at (**). By assumption (2), we know that $(\mathbf{x}_i, u_i)$ is i.i.d. Furthermore, $\mathbb{E}[\mathbf{x}_i u_i] = \mathbf{0}$, so the first moment exists. We can therefore apply the **weak law of large numbers** to get:

$$\frac{\mathbf{X}^T \mathbf{u}}{n} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i u_i \xrightarrow{P} \mathbb{E}[\mathbf{x}_i u_i] =_{(1)} \mathbf{0}$$

whereby $=_{(1)}$ arises from the orthogonal assumption (3) of $\mathbf{X}$ and $\mathbf{u}$.

Therefore, we can apply the **continuous mapping theorem** to the product of (*) and (**) that we started with, whereby we now have that the **product of the limits is the limit of the products**:

$$p \lim_{n \rightarrow \infty} \left[\left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1} \left(\frac{\mathbf{X}^T \mathbf{u}}{n} \right) \right] = \left[p \lim_{n \rightarrow \infty} \left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1} \right] \left[p \lim_{n \rightarrow \infty} \left(\frac{\mathbf{X}^T \mathbf{u}}{n} \right) \right] = \left[M_{XX}^{-1} \right] \left[\mathbf{0} \right] = \mathbf{0}$$

We can therefore conclude that:

$$p \lim_{n \rightarrow \infty} \hat{\boldsymbol{\beta}}_{OLS} = p \lim_{n \rightarrow \infty} \boldsymbol{\beta} + p \lim_{n \rightarrow \infty} \left[\left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1} \left(\frac{\mathbf{X}^T \mathbf{u}}{n} \right) \right] = \boldsymbol{\beta} + \mathbf{0}$$

Therefore, we have shown consistency:

$$p \lim_{n \rightarrow \infty} \hat{\boldsymbol{\beta}}_{OLS} = \boldsymbol{\beta}.$$

□

In our proof, we assumed that the data $(y_i, \mathbf{x}_i^T)$ were independent and identically distributed. This is stronger than necessary: consistency can also be established for independent but non-identically distributed observations when suitable laws of large numbers apply to $n^{-1} \sum_i \mathbf{x}_i \mathbf{x}_i^T$ and $n^{-1} \sum_i \mathbf{x}_i u_i$.

Theorem 217 (OLS Consistency under Weaker Assumptions)

Suppose that the data $(y_i, \mathbf{x}_i^T)$ are independent over $i = 1, 2, \dots, n$ but **not** identically distributed. If $n^{-1} \sum_i \mathbf{x}_i \mathbf{x}_i^T$ converges in probability to a finite non-singular matrix and $n^{-1} \sum_i \mathbf{x}_i u_i \xrightarrow{P} \mathbf{0}$ under suitable moment conditions, then the OLS estimator is still **consistent**

$$p \lim_{n \rightarrow \infty} \hat{\boldsymbol{\beta}}_{OLS} = \boldsymbol{\beta} \tag{166}$$

or equivalently:

$$\hat{\boldsymbol{\beta}}_{OLS} \xrightarrow{P} \boldsymbol{\beta}. \tag{167}$$

Proof. See section 4.4 of Cameron and Trivedi (2005).

□

7.7.2 Linear Projection Model Revisited

So far, the key assumption we required was the **orthogonality assumption**

$$\mathbb{E}[\mathbf{x}_i u_i] = 0.$$

In practice, we generally consider linear models whereby:

1. We have an intercept term so that $x_{1i} = 1$ for all $i = 1, 2, \dots, n$. This then implies that the first element of the $K \times 1$ vector $\mathbb{E}[\mathbf{x}_i u_i] = 0$ implies that

$$\mathbb{E}[x_{1i} u_i] = \mathbb{E}[(1) u_i] = \mathbb{E}[u_i] = 0.$$

2. Alternatively, we can assume that $\mathbb{E}[y_i] = 0$ and $\mathbb{E}[\mathbf{x}_i] = \mathbf{0}$. This then implies that:

$$\mathbb{E}[u_i] = 0.$$

Since in both cases, we have that $\mathbb{E}[u_i] = 0$ and the orthogonality assumption $\mathbb{E}[\mathbf{x}_i u_i] = 0$, this then implies that:

$$\text{Cov}[\mathbf{x}_i, u_i] = \mathbb{E}[\mathbf{x}_i u_i] - \mathbb{E}[\mathbf{x}_i] \mathbb{E}[u_i] = 0$$

and hence:

$$\text{Corr}[\mathbf{x}_i, u_i] = 0.$$

This allows us to define the **linear projection model** based off this orthogonality assumption.

Definition 218 (Linear Projection Model)

We make the assumptions that:

$$y_i = \mathbf{x}_i^T \boldsymbol{\beta} + u_i$$

and

$$\mathbb{E}[\mathbf{x}_i u_i] = \mathbf{0}.$$

The **linear projection** model is then defined as:

$$\mathbb{E}^*[y_i | \mathbf{x}_i] = \mathbf{x}_i^T \boldsymbol{\beta}$$

Remark 219. The linear projection model examines the linear relationship between y_i and the explanatory variable $\mathbf{x}_i$ with the property that the additive error term u_i is **uncorrelated in the population equation** with each of the included explanatory variables in $\mathbf{x}_i$.

In our traditional linear regression model setting, we made the assumption of **linear conditional expectation**:

$$\mathbb{E}[y_i | \mathbf{x}_i] = \mathbf{x}_i^T \boldsymbol{\beta} \quad \leftrightarrow \quad \mathbb{E}[u_i | \mathbf{x}_i] = 0.$$

Definition 220 (Linear Regression Model)

We make the assumptions that:

$$\mathbb{E}[y_i | \mathbf{x}_i] = \mathbf{x}_i^T \boldsymbol{\beta} \quad \leftrightarrow \quad \mathbb{E}[u_i | \mathbf{x}_i] = 0.$$

The **linear regression** model is then defined as:

$$y_i = \mathbf{x}_i^T \boldsymbol{\beta} + u_i$$

This in fact implies the orthogonality assumption in the linear projection model.

Proposition 221 (Linear Conditional Expectation Implies Orthogonality)

The linear conditional expectation assumption implies the orthogonality assumption:

$$\mathbb{E}[y_i | \mathbf{x}_i] = \mathbf{x}_i^T \boldsymbol{\beta} \quad \Rightarrow \quad \mathbb{E}[\mathbf{x}_i u_i] = 0$$

Proof. Apply the law of iterated expectations:

$$\begin{aligned} \mathbb{E}[\mathbf{x}_i u_i] &= \mathbb{E}_{\mathbf{x}}[\mathbb{E}_{u|\mathbf{x}}[\mathbf{x}_i u_i | \mathbf{x}_i]] \\ &= \mathbb{E}_{\mathbf{x}}[\mathbf{x}_i \mathbb{E}_{u|\mathbf{x}}[u_i | \mathbf{x}_i]] \\ &=_{(1)} \mathbb{E}_{\mathbf{x}}[0] \\ &= 0 \end{aligned}$$

whereby $=_{(1)}$ comes from the linear conditional expectation assumption $\mathbb{E}[u_i | \mathbf{x}_i] = 0$. □

Theorem 222 (Linear Regression Model Implies Linear Projection Model)

The linear regression model is a *model subset* of the linear projection model.

It is not the case that the converse holds i.e. $\mathbb{E}[\mathbf{x}_i u_i] = 0$ **does not** imply that $\mathbb{E}[u_i | \mathbf{x}_i] = 0$. It is a weaker assumption as we are not specifying that the conditional expectation $\mathbb{E}[y_i | \mathbf{x}_i]$ be a linear function.

Furthermore, the orthogonality function for the linear regression model, whereby for any bounded function $g(\mathbf{x}_i)$, we have that:

$$\mathbb{E}\left[\{y_i - \mathbb{E}[y_i | \mathbf{x}_i]\}g(\mathbf{x}_i)\right] = \mathbb{E}\left[\{y_i - \mathbf{x}_i^T \boldsymbol{\beta}\}g(\mathbf{x}_i)\right] = 0$$

implies the orthogonality condition in the linear projection model (whereby $g(\mathbf{x}_i) = \mathbf{x}_i$):

$$\mathbb{E}\left[\{y_i - \mathbf{x}_i^T \boldsymbol{\beta}\}\mathbf{x}_i\right] = \mathbb{E}\left[u_i \mathbf{x}_i\right] = 0.$$

However, the converse **does not hold**.

7.7.3 Asymptotic Normality

We now want to show the limiting distribution of the OLS estimator where we do not have the assumption of normality of the error terms. Doing so will allow us to construct confidence intervals and conduct asymptotically valid hypothesis testing.

Assumption 223 (Assumptions for Asymptotic Normality)

The assumptions needed for asymptotic normality of the OLS estimator are:

1. $\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}$;
2. The data on $(y_i, \mathbf{x}_i^T)$ are independent and identically distributed for $i = 1, \dots, n$;
3. $\mathbb{E}[\mathbf{x}_i u_i] = 0$ for all $i = 1, \dots, n$;
4. The $K \times K$ matrix $M_{XX} = \mathbb{E}[\mathbf{x}_i \mathbf{x}_i^T]$ exists and is non-singular;
5. The $K \times K$ matrix $M_{X\Omega X} = \mathbb{E}[u_i^2 \mathbf{x}_i \mathbf{x}_i^T]$ exists and is non-singular;

We now state the most important result of this section.

Theorem 224 (Asymptotic Normality of OLS Estimator)

Under the assumptions for asymptotic normality, we have that:

$$\sqrt{n}[\hat{\boldsymbol{\beta}}_{OLS} - \boldsymbol{\beta}] \xrightarrow{D} \mathcal{N}(\mathbf{0}, M_{XX}^{-1} M_{X\Omega X} M_{XX}^{-1}). \quad (168)$$

Proof. First, recall that the expression for the OLS estimator is:

$$\hat{\boldsymbol{\beta}}_{OLS} = \boldsymbol{\beta} + \left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1} \left(\frac{\mathbf{X}^T \mathbf{u}}{n} \right)$$

or alternatively:

$$\hat{\boldsymbol{\beta}}_{OLS} - \boldsymbol{\beta} = \left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1} \left(\frac{\mathbf{X}^T \mathbf{u}}{n} \right)$$

We then rescale it by $\sqrt{n}$ to get:

$$\sqrt{n}[\hat{\boldsymbol{\beta}}_{OLS} - \boldsymbol{\beta}] = \underbrace{\left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1}}_{(*)} \underbrace{\left(\frac{\mathbf{X}^T \mathbf{u}}{\sqrt{n}} \right)}_{(**)}$$

We look at the limiting result for $(*)$ and $(**)$. First, for $(*)$, we can apply the **weak law of large numbers**:

$$\frac{\mathbf{X}^T \mathbf{X}}{n} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^T \xrightarrow{P} \mathbb{E}[\mathbf{x}_i \mathbf{x}_i^T] =_{(1)} M_{XX}$$

whereby M_{XX} holds by definition of M_{XX} in assumption (4). Furthermore, we know that by assumption (4), the inverse of M_{XX} exists, so by the **continuous mapping theorem**, we have that

$$\left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1} \xrightarrow{P} (M_{XX})^{-1} = M_{XX}^{-1}.$$

Now looking at $(**)$, we can rewrite $(**)$ as:

$$\frac{\mathbf{X}^T \mathbf{u}}{\sqrt{n}} = \sqrt{n} \left(\frac{\mathbf{X}^T \mathbf{u}}{n} \right) = \sqrt{n} \left(\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i u_i \right)$$

where we now have a $K \times 1$ vector of sample means: $\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i u_i$. This gives us a hint to apply the **central limit theorem** to this vector of sample means. First, we compute the mean and variance of this expression:

$$\mathbb{E}[\mathbf{x}_i u_i] =_{(1)} \mathbf{0}$$

whereby $=_{(1)}$ comes from assumption (3) of orthogonality. For the variance, we have:

$$\begin{aligned} \text{Var}[\mathbf{x}_i u_i] &= \mathbb{E}[(\mathbf{x}_i u_i)(\mathbf{x}_i u_i)^T] - \mathbb{E}[\mathbf{x}_i u_i] \mathbb{E}[\mathbf{x}_i u_i]^T \\ &= \mathbb{E}[u_i^2 \mathbf{x}_i \mathbf{x}_i^T] - \mathbf{0} \\ &=_{(1)} M_{X\Omega X} \end{aligned}$$

whereby $=_{(1)}$ comes from the definition of $M_{X\Omega X}$ in assumption (5). Finally, the $K \times 1$ random vectors $(\mathbf{x}_i u_i)$ are independent and identically distributed over $i = 1, 2, \dots, n$ by assumption (2). We can therefore apply the central limit theorem to (**):

$$\frac{\mathbf{X}^T \mathbf{u}}{\sqrt{n}} = \sqrt{n} \left(\frac{\mathbf{X}^T \mathbf{u}}{n} \right) \xrightarrow{D} \mathcal{N}(\mathbf{0}, M_{X\Omega X}).$$

We now go back to our main expression:

$$\sqrt{n}[\hat{\boldsymbol{\beta}}_{OLS} - \boldsymbol{\beta}] = \underbrace{\left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1}}_{(*)} \underbrace{\left(\frac{\mathbf{X}^T \mathbf{u}}{\sqrt{n}} \right)}_{(**)}$$

We have shown that $(*) \xrightarrow{P} M_{XX}^{-1}$ and $(**) \xrightarrow{D} \mathcal{N}(\mathbf{0}, M_{X\Omega X})$. We apply **Slutsky's theorem** to conclude that:

$$\begin{aligned} \left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1} \left(\frac{\mathbf{X}^T \mathbf{u}}{\sqrt{n}} \right) &\xrightarrow{D} M_{XX}^{-1} \cdot \mathcal{N}(\mathbf{0}, M_{X\Omega X}) = \mathcal{N}(M_{XX}^{-1} \cdot \mathbf{0}, M_{XX}^{-1} M_{X\Omega X} (M_{XX}^{-1})^T) \\ &=_{(1)} \mathcal{N}(\mathbf{0}, M_{XX}^{-1} M_{X\Omega X} M_{XX}^{-1}) \end{aligned}$$

whereby $=_{(1)}$ uses the fact that M_{XX}^{-1} is **symmetric**. Therefore, we can conclude our desired result:

$$\sqrt{n}[\hat{\boldsymbol{\beta}}_{OLS} - \boldsymbol{\beta}] \xrightarrow{D} \mathcal{N}(\mathbf{0}, M_{XX}^{-1} M_{X\Omega X} M_{XX}^{-1}).$$

□

Remark 225. We could have relaxed the assumption that $(y_i, \mathbf{x}_i^T)$ is i.i.d. to independence under suitable non-identically distributed central limit theorem conditions.

This brings us to a broader and more general definition.

Definition 226 (Asymptotically Normal)

An estimator $\hat{\boldsymbol{\theta}}$ of the parameter vector $\boldsymbol{\theta}$ with the property that:

$$\sqrt{n}(\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}) \xrightarrow{D} \mathcal{N}(\mathbf{0}, V) \tag{169}$$

for some finite variance matrix V is said to be **asymptotically normal**.

We are saying that if appropriately scaled, the estimator has a non-degenerate limit distribution which is normal. Using this terminology, we say that the OLS estimator is **asymptotically normal**.

We now wish to explore further properties of asymptotic normality. We now look at a *special case* of the classical linear regression model with linear conditional expectation and conditional homoskedasticity.

Assumption 227 (Assumptions for Asymptotic Normality)

The assumptions needed for asymptotic normality of the OLS estimator are:

1. $\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}$;
2. The data on $(y_i, \mathbf{x}_i^T)$ are independent and identically distributed for $i = 1, \dots, n$;
3. Linear conditional expectation $\mathbb{E}[u_i | \mathbf{x}_i] = 0$ for all $i = 1, \dots, n$;
4. The $K \times K$ matrix $M_{XX} = \mathbb{E}[\mathbf{x}_i \mathbf{x}_i^T]$ exists and is non-singular;
5. The $K \times K$ matrix $M_{X\Omega X} = \mathbb{E}[u_i^2 \mathbf{x}_i \mathbf{x}_i^T]$ exists and is non-singular;
6. Conditional homoskedasticity $\text{Var}[u_i | \mathbf{x}_i] = \mathbb{E}[u_i^2 | \mathbf{x}_i] = \sigma^2$.

We now have a different limiting variance matrix for our OLS estimator.

Theorem 228 (Asymptotic Normality of OLS Estimator under Homoskedasticity)

Under the strengthened assumptions for asymptotic normality with conditional homoskedasticity, we have that:

$$\sqrt{n}[\hat{\boldsymbol{\beta}}_{OLS} - \boldsymbol{\beta}] \xrightarrow{D} \mathcal{N}(\mathbf{0}, \sigma^2 M_{XX}^{-1}). \quad (170)$$

Proof. Under assumption (6), we have that:

$$\begin{aligned} M_{X\Omega X} &= \mathbb{E}[u_i^2 \mathbf{x}_i \mathbf{x}_i^T] \\ &=_{(1)} \mathbb{E}_{\mathbf{X}} \left[\mathbb{E}_{u|\mathbf{X}}[u_i^2 \mathbf{x}_i \mathbf{x}_i^T | \mathbf{x}_i] \right] \\ &= \mathbb{E}_{\mathbf{X}} \left[\mathbf{x}_i \mathbf{x}_i^T \mathbb{E}_{u|\mathbf{X}}[u_i^2 | \mathbf{x}_i] \right] \\ &=_{(2)} \mathbb{E}_{\mathbf{X}} \left[\mathbf{x}_i \mathbf{x}_i^T \sigma^2 \right] \\ &= \sigma^2 \mathbb{E}_{\mathbf{X}} \left[\mathbf{x}_i \mathbf{x}_i^T \right] \\ &= \sigma^2 M_{XX} \end{aligned}$$

where $=_{(1)}$ is the application of the law of iterated expectations and $=_{(2)}$ is from assumption (6) of conditional homoskedasticity.

We can therefore apply this to our original limiting distribution result:

$$\sqrt{n}[\hat{\boldsymbol{\beta}}_{OLS} - \boldsymbol{\beta}] \xrightarrow{D} \mathcal{N}(\mathbf{0}, M_{XX}^{-1} M_{X\Omega X} M_{XX}^{-1}) = \mathcal{N}(\mathbf{0}, M_{XX}^{-1} (\sigma^2 M_{XX}) M_{XX}^{-1}) = \mathcal{N}(\mathbf{0}, \sigma^2 M_{XX}^{-1})$$

□

We therefore have that nice result that:

$$\sqrt{n}[\hat{\boldsymbol{\beta}}_{OLS} - \boldsymbol{\beta}] \xrightarrow{D} \mathcal{N}(\mathbf{0}, V)$$

where the limiting variance $V = \sigma^2 M_{XX}^{-1}$. We can rewrite this expression for the OLS estimator to be asymptotically normal:

$$\hat{\boldsymbol{\beta}}_{OLS} \overset{a}{\sim} \mathcal{N}\left(\boldsymbol{\beta}, \frac{V}{n}\right)$$

where $\overset{a}{\sim}$ denotes the asymptotic distribution. This is called the **asymptotic distribution** of $\hat{\beta}_{OLS}$. However, this is not practical to use. To see why, look at the asymptotic variance:

$$V = \sigma^2 M_{XX}^{-1}.$$

We have shown that:

$$\left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1} \xrightarrow{P} M_{XX}^{-1}$$

so that means we have a **consistent estimator** of the $K \times K$ matrix M_{XX}^{-1} . However, we do not know what σ^2 is. Therefore, we need an estimator for σ^2 .

We can use either the OLS estimator:

$$\hat{\sigma}_{OLS}^2 = \frac{1}{n-K} \sum_{i=1}^n \hat{u}_i^2 = \frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{n-K} \quad (171)$$

or the MLE estimator:

$$\hat{\sigma}_{MLE}^2 = \frac{1}{n} \sum_{i=1}^n \hat{u}_i^2 = \frac{\hat{\mathbf{u}}^T \hat{\mathbf{u}}}{n} \quad (172)$$

to estimate the unknown variance σ^2 . Both of these estimators are consistent. Therefore, we can apply **Slutsky's theorem** on the estimator:

$$\hat{V} = \hat{\sigma}_{OLS}^2 \left(\frac{\mathbf{X}^T \mathbf{X}}{n} \right)^{-1} \xrightarrow{P} V$$

which means we have a consistent estimator:

$$\hat{V} \xrightarrow{P} V$$

We are able to replace V with $\hat{V}$ in our asymptotic distribution without changing any of our asymptotic distribution results. We therefore have the approximation of the limiting distribution:

$$\hat{\beta}_{OLS} \overset{a}{\sim} \mathcal{N}(\beta, \frac{\hat{V}}{n})$$

which we can use to construct approximate confidence intervals and asymptotically valid test statistics. Note the limiting variance and the asymptotic distribution for $\hat{\beta}_{OLS}$ coincides with the **exact finite** sample conditional distribution of $\hat{\beta}_{OLS}|\mathbf{X}$ we made when we assumed that the error terms were normally distributed. Hence, we know that in this scenario, our asymptotic approximation will be accurate.

7.7.4 Asymptotic Inference

We state the test statistic we can use for **simple hypothesis testing**.

Proposition 229 (Asymptotic T-test Statistic For Simple Hypothesis Testing)

We can use the following test statistic for **simple hypothesis testing** of the form $H_0 : \beta_k = \beta_k^0$:

$$t_k = \frac{\hat{\beta}_k - \beta_k^0}{se_k} \overset{a}{\sim} \mathcal{N}(0, 1) \quad \text{under } H_0. \quad (173)$$

where se_k is the asymptotic standard error.

Proof. Recall our asymptotic distribution:

$$\hat{\beta}_{OLS} \overset{a}{\sim} \mathcal{N}(\beta, \frac{\hat{V}}{n})$$

Let $\hat{v}_{kk}$ denote the k^{th} element of the main diagonal of the $K \times K$ variance matrix estimator:

$$\frac{\hat{V}}{n} = \hat{\sigma}_{OLS}^2 (\mathbf{X}^T \mathbf{X})^{-1}.$$

We have that for each OLS estimator β_k :

$$\hat{\beta}_k \stackrel{a}{\sim} \mathcal{N}(\beta_k, \hat{v}_{kk})$$

We can standardise this to get:

$$t_k = \frac{\hat{\beta}_k - \beta_k}{\sqrt{\hat{v}_{kk}}} = \frac{\hat{\beta}_k - \beta_k}{se_k} \stackrel{a}{\sim} \mathcal{N}(0, 1).$$

□

We can then construct **simple hypothesis tests** of the form $H_0 : \beta_k = \beta_k^0$ by comparing the test statistic

$$t_k = \frac{\hat{\beta}_k - \beta_k^0}{se_k}$$

computed under the null hypothesis and comparing it with the standard normal distribution $\mathcal{N}(0, 1)$ at the desired significance level. We can also justify the asymptotic test statistic we just derived by looking at the limiting distribution of a finite sample t-test statistic. It may seem strange that we were able to treat t_k as being drawn from the standard normal distribution even though we estimated the unknown σ^2 . We can in fact show that this is entirely justified.

Theorem 230 (Asymptotic Validity of T-test Statistic)

Define the z-statistic as:

$$z_k = \frac{\hat{\beta}_k - \beta_k}{\sqrt{v_{kk}}} \xrightarrow{D} \mathcal{N}(0, 1) \quad (174)$$

We can also show that the t-statistic with an estimated v_{kk} is asymptotically equivalent:

$$t_k = \frac{\hat{\beta}_k - \beta_k}{\sqrt{\hat{v}_{kk}}} \xrightarrow{D} \mathcal{N}(0, 1) \quad (175)$$

Proof. Under the asymptotic normality result for OLS:

$$\sqrt{n}(\hat{\beta}_{OLS} - \beta) \xrightarrow{D} \mathcal{N}(\mathbf{0}, \sigma^2 M_{XX}^{-1}).$$

This automatically implies that the standardised statistic:

$$z_k = \frac{\hat{\beta}_k - \beta_k}{\sqrt{v_{kk}}} \xrightarrow{D} \mathcal{N}(0, 1).$$

We now look at the t-statistic and rewrite it as:

$$t_k = \frac{\hat{\beta}_k - \beta_k}{\sqrt{\hat{v}_{kk}}} = \frac{\sqrt{v_{kk}}}{\sqrt{\hat{v}_{kk}}} \left(\frac{\hat{\beta}_k - \beta_k}{\sqrt{v_{kk}}} \right) = \frac{\sqrt{v_{kk}}}{\sqrt{\hat{v}_{kk}}} z_k$$

We have already shown that $z_k \xrightarrow{D} \mathcal{N}(0, 1)$. We look at the term $\frac{\sqrt{v_{kk}}}{\sqrt{\hat{v}_{kk}}}$, which contains the true variance in the numerator and the estimated variance in the denominator. Since $\hat{\sigma}^2/\sigma^2 \xrightarrow{P} 1$, the corresponding variance estimator is ratio-consistent:

$$\frac{\hat{v}_{kk}}{v_{kk}} \xrightarrow{P} 1.$$

We can then apply the **continuous mapping theorem**:

$$\frac{\sqrt{\hat{V}_{kk}}}{\sqrt{V_{kk}}} \xrightarrow{P} 1.$$

We then apply **Slutsky's theorem**:

$$\frac{\sqrt{V_{kk}}}{\sqrt{\hat{V}_{kk}}} \xrightarrow{P} 1.$$

We conclude and show our result by again applying Slutsky's theorem:

$$t_k = \frac{\sqrt{V_{kk}}}{\sqrt{\hat{V}_{kk}}} z_k \xrightarrow{D} (1) \cdot \mathcal{N}(0, 1) = \mathcal{N}(0, 1).$$

□

We now want to test linear restrictions $\theta = H\beta$ whereby H is a non-stochastic $p \times K$ matrix with $\text{rank}(H) = p$ (full rank). Go back to our asymptotic distribution for the OLS estimator:

$$\hat{\beta}_{OLS} \overset{a}{\sim} \mathcal{N}\left(\beta, \frac{\hat{V}}{n}\right)$$

The asymptotic distribution of the $p \times 1$ random vector $\hat{\theta}_{OLS} = H\hat{\beta}_{OLS}$ is:

$$\hat{\theta}_{OLS} \overset{a}{\sim} \mathcal{N}\left(H\beta, H\frac{\hat{V}}{n}H^T\right)$$

Proposition 231 (Consistency of Linear Restrictions)

The estimator of θ constructed from the OLS estimator $\hat{\theta}_{OLS} = H\hat{\beta}_{OLS}$ is consistent:

$$p \lim_{n \rightarrow \infty} \hat{\theta}_{OLS} = \theta.$$

Proof. By consistency of $\hat{\beta}_{OLS}$, we have the consistency result:

$$p \lim_{n \rightarrow \infty} \hat{\theta}_{OLS} = p \lim_{n \rightarrow \infty} (H\hat{\beta}_{OLS}) = H(p \lim_{n \rightarrow \infty} \hat{\beta}_{OLS}) = H\beta = \theta.$$

□

We can then construct the scalar random variable w as our test statistic:

$$w = (\hat{\theta}_{OLS} - \theta^0)^T \left[H\frac{\hat{V}}{n}H^T \right]^{-1} (\hat{\theta}_{OLS} - \theta^0)$$

This is similar to what we did earlier when we were constructing an F-test in finite sample. We can then show the asymptotic property of this test statistic.

Proposition 232 (Asymptotic Distribution of Quadratic Test Statistic)

The asymptotic distribution of the quadratic form test statistic for testing p linear restrictions is given by:

$$w = (\hat{\theta}_{OLS} - \theta^0)^T \left[H\frac{\hat{V}}{n}H^T \right]^{-1} (\hat{\theta}_{OLS} - \theta^0) \overset{a}{\sim} \chi^2(p) \quad \text{under } H_0. \quad (176)$$

Proof. The proof of this distribution follows the proof for the asymptotic distribution of the F-test.

□

Under the special case of the classical linear regression model, we have that:

$$\hat{V}/n = \hat{\sigma}_{OLS}^2 (\mathbf{X}^T \mathbf{X})^{-1}.$$

Therefore, we have that under conditional homoskedasticity:

$$w = (\hat{\theta}_{OLS} - \theta^0)^T \left[\hat{\sigma}_{OLS}^2 H(\mathbf{X}^T \mathbf{X})^{-1} H^T \right]^{-1} (\hat{\theta}_{OLS} - \theta^0) \stackrel{a}{\sim} \chi^2(p) \quad \text{under } H_0.$$

This has an important name.

Definition 233 ((Asymptotic) Wald Statistic)

The Wald statistic is given by:

$$w = (\hat{\theta}_{OLS} - \theta^0)^T \left[\hat{\sigma}_{OLS}^2 H(\mathbf{X}^T \mathbf{X})^{-1} H^T \right]^{-1} (\hat{\theta}_{OLS} - \theta^0) \stackrel{a}{\sim} \chi^2(p) \quad \text{under } H_0. \quad (177)$$

Therefore, joint hypothesis tests of the form: $H_0 : H\beta = \theta^0$ can be conducted by computing the test statistic:

$$w = (\hat{\theta}_{OLS} - \theta^0)^T \left[\hat{\sigma}_{OLS}^2 H(\mathbf{X}^T \mathbf{X})^{-1} H^T \right]^{-1} (\hat{\theta}_{OLS} - \theta^0)$$

which has a distribution that is well approximated in large finite samples by the $\chi^2(p)$ distribution if the null hypothesis is true. We can compare the value of the test statistic to the relevant critical values of the $\chi^2(p)$ distribution at the desired level of significance. This is known as the **Wald test**.

Alternatively, we can justify this χ^2 distribution by another asymptotic argument. We can take our (finite sample) F-test for linear restrictions and show that it converges to a χ^2 distribution.

Proposition 234 (Asymptotic Distribution of F distribution is chi-squared distribution)

Denote $v \sim F(m, n)$. Then

$$mv \xrightarrow{D} \chi^2(m) \quad (178)$$

as $n \rightarrow \infty$.

Proof. First, we write out the test statistic:

$$v = \frac{w_1/m}{w_2/n}$$

whereby $w_1 \sim \chi^2(m)$ and $w_2 \sim \chi^2(n)$. This then implies that:

$$mv = \frac{w_1}{w_2/n}.$$

We now look at the denominator. By definition of χ^2 -distribution, we have that we can express w_2 as the sum of independent standard normals:

$$w_2 = \sum_{i=1}^n z_i^2$$

whereby $z_i \sim \mathcal{N}(0, 1)$ for $i = 1, 2, \dots, n$ are independent standard normal random variables. First, we note that

$$\mathbb{E}[z_i] = 0$$

and

$$\text{Var}[z_i] = \mathbb{E}[z_i^2] - \mathbb{E}[z_i]^2 \leftrightarrow 1 = \mathbb{E}[z_i^2] - 0^2$$

$$\Rightarrow \mathbb{E}[z_i^2] = 1$$

We can see therefore see that:

$$\begin{aligned}\mathbb{E}[w_2] &= \mathbb{E}\left[\sum_{i=1}^n z_i^2\right] \\ &= n\mathbb{E}[z_i^2] \\ &=_{(1)} n \times 1 \\ &= n.\end{aligned}$$

Hence,

$$\mathbb{E}\left[\frac{w_2}{n}\right] = \frac{n}{n} = 1.$$

We have verified that z_i are independent and has finite first and second moments. We can therefore apply the **weak law of large numbers** to the denominator:

$$\frac{w_2}{n} = \frac{1}{n}w_2 = \frac{1}{n}\sum_{i=1}^n z_i^2 \xrightarrow{P} \mathbb{E}[z_i^2] = 1.$$

We can therefore apply both the weak law of large numbers and **Slutsky's theorem** to our original test statistic to get our desired result:

$$mv = \frac{w_1}{w_2/n} \xrightarrow{D} w_1 \sim \chi^2(m).$$

□

Hence, in large samples, to test p linear restrictions, we can use the test statistic pv and critical values obtained from the $\chi^2(p)$ distribution.

8 Heteroskedasticity and Generalised Least Squares

8.1 Conditional Heteroskedasticity

We now drop the assumption of homoskedasticity whereby the conditional variance

$$\text{Var}[u_i|\mathbf{x}_i] = \sigma^2 \quad \text{for } i = 1, \dots, n$$

is common to all observations.

Assumption 235 (Conditional Heteroskedasticity)

The conditional variance is allowed to be different for different observations:

$$\text{Var}[u_i|\mathbf{x}_i] = \sigma_i^2 \quad \text{for } i = 1, \dots, n.$$

It is important to note that neither consistency nor asymptotic normality is lost when we allow for conditional heteroskedasticity, provided the other assumptions continue to hold. Instead, the variance of the scaled estimator will be different.

Assumption 236 (Assumptions for Asymptotic Normality)

The assumptions needed for asymptotic normality of the OLS estimator are:

1. $\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}$;
2. The data on $(y_i, \mathbf{x}_i^T)$ are independent and identically distributed for $i = 1, \dots, n$;
3. $\mathbb{E}[\mathbf{x}_i u_i] = 0$ for all $i = 1, \dots, n$;
4. The $K \times K$ matrix $M_{XX} = \mathbb{E}[\mathbf{x}_i \mathbf{x}_i^T]$ exists and is non-singular;
5. The $K \times K$ matrix $M_{X\Omega X} = \mathbb{E}[u_i^2 \mathbf{x}_i \mathbf{x}_i^T]$ exists and is non-singular;

Recall that under these assumptions, we have the asymptotic normality of the OLS estimator:

$$\sqrt{n}[\hat{\boldsymbol{\beta}}_{OLS} - \boldsymbol{\beta}] \xrightarrow{D} \mathcal{N}(\mathbf{0}, M_{XX}^{-1} M_{X\Omega X} M_{XX}^{-1}). \quad (179)$$

Alternatively, we can write this as:

$$\hat{\boldsymbol{\beta}}_{OLS} \overset{a}{\sim} \mathcal{N}(\boldsymbol{\beta}, \frac{\mathbf{V}}{n})$$

whereby the variance-covariance matrix is:

$$\frac{\mathbf{V}}{n} = \left(\frac{1}{n}\right) M_{XX}^{-1} M_{X\Omega X} M_{XX}^{-1}$$

We now want to find an estimator $\hat{\mathbf{V}}$ for $\mathbf{V}$ under conditional heteroskedasticity. We state what this estimator is.

Theorem 237 (Heteroskedasticity-Consistent Estimator for Variance Matrix)

The following estimator $\hat{\mathbf{V}}$ is a **consistent** estimator for the variance-covariance matrix $\mathbf{V}$:

$$\hat{\mathbf{V}} = \left(\frac{\mathbf{X}^T \mathbf{X}}{n}\right)^{-1} \left(\frac{1}{n} \sum_{i=1}^n \hat{u}_i^2 \mathbf{x}_i \mathbf{x}_i^T\right) \left(\frac{\mathbf{X}^T \mathbf{X}}{n}\right)^{-1} \xrightarrow{P} M_{XX}^{-1} M_{X\Omega X} M_{XX}^{-1} = \mathbf{V} \quad (180)$$

Proof. We have previously shown that the *bread* in our estimator:

$$\left(\frac{\mathbf{X}^T \mathbf{X}}{n}\right)^{-1} \xrightarrow{P} M_{XX}^{-1}$$

gave us a consistent estimator of the $K \times K$ matrix M_{XX}^{-1} . We now want to find an estimator for the $K \times K$ matrix $M_{X\Omega X}$.

Theorem 238 (Consistent Estimator for Filling)

The following estimator is a consistent estimator for the filling:

$$\frac{1}{n} \sum_{i=1}^n \hat{u}_i^2 \mathbf{x}_i \mathbf{x}_i^T \xrightarrow{P} \mathbb{E}[u_i^2 \mathbf{x}_i \mathbf{x}_i^T] = M_{X\Omega X} \quad (181)$$

whereby $\hat{u}_i \xrightarrow{P} u_i$.

We can then use Slutsky's theorem to get our desired result:

$$\hat{\mathbf{V}} = \left(\frac{\mathbf{X}^T \mathbf{X}}{n}\right)^{-1} \left(\frac{1}{n} \sum_{i=1}^n \hat{u}_i^2 \mathbf{x}_i \mathbf{x}_i^T\right) \left(\frac{\mathbf{X}^T \mathbf{X}}{n}\right)^{-1} \xrightarrow{P} M_{XX}^{-1} M_{X\Omega X} M_{XX}^{-1} = \mathbf{V}$$

□

Therefore, as $\hat{\mathbf{V}} \xrightarrow{P} \mathbf{V}$, then we can replace $\mathbf{V}$ by $\hat{\mathbf{V}}$ without affecting the limiting distribution result:

$$\hat{\boldsymbol{\beta}}_{OLS} \overset{a}{\sim} \mathcal{N}\left(\boldsymbol{\beta}, \frac{\hat{\mathbf{V}}}{n}\right)$$

whereby:

$$\frac{\hat{\mathbf{V}}}{n} = (\mathbf{X}^T \mathbf{X})^{-1} \left(\sum_{i=1}^n \hat{u}_i^2 \mathbf{x}_i \mathbf{x}_i^T\right) (\mathbf{X}^T \mathbf{X})^{-1}$$

is the heteroskedastic-consistent estimator whereby we can compute $\frac{\hat{\mathbf{V}}}{n}$ using the data on $\mathbf{X}$ and the OLS residuals $\hat{\mathbf{u}}$. We can then conduct hypothesis testing with the standard errors that arise from this estimator.

Definition 239 (White Standard Errors)

The White Standard Errors (heteroskedasticity-robust standard errors) are the square roots of the scalar elements on the main diagonal of $\hat{\mathbf{V}}/n$.

Note that conditional homoskedasticity $\mathbb{E}[u_i^2 | \mathbf{x}] = \sigma^2$ is a *special case* of conditional heteroskedasticity and hence we can use the White standard errors even in the case of homoskedasticity.

We can also conduct tests for conditional heteroskedasticity based on the OLS residuals.

Proposition 240 (White Test)

The White test for conditional heteroskedasticity is found by regressing the squared OLS residuals $\hat{u}_i^2$ on a constant and on all the explanatory variables and their squares and their cross-products.

That is, first retrieve the residuals $\hat{u}_i$ from the regression:

$$y_i = \beta_1 + \beta_2 x_{2i} + \dots \beta_K x_{Ki} + u_i$$

and then run the auxiliary regression:

$$\hat{u}_i^2 = \gamma_1 + \mathbf{z}_i^T \boldsymbol{\gamma} + v_i,$$

where $\mathbf{z}_i$ contains all non-redundant explanatory variables, their squares, and their cross-products. Then jointly test the restriction:

$$H_0 : \boldsymbol{\gamma} = \mathbf{0}.$$

If we fail to reject the null, then we have insufficient evidence of conditional heteroskedasticity; failure to reject does not prove homoskedasticity.

The intuition behind the test is to see whether is the conditional variance a function of the observed variables $\mathbf{x}_i$.

Under the assumptions stated above, the OLS estimator remains consistent and asymptotically normal under heteroskedasticity. White standard errors consistently estimate its asymptotic variance and therefore make asymptotic inference valid. However, this is a passive response to heteroskedasticity as we are not specifying its form and hence, the OLS estimator is generally **not** asymptotically efficient.

8.2 Generalised Least Squares

We can construct a more efficient estimator if we are willing to specify what is the form of the conditional heteroskedasticity. The process of generalised least squares does this. We need to impose some assumptions before we can proceed.

Assumption 241 (Generalised Least Squares Assumptions)

We assume that:

1. $\mathbb{E}[\mathbf{y}|\mathbf{X}] = \mathbf{X}\boldsymbol{\beta}$
2. $\text{Var}[\mathbf{y}|\mathbf{X}] = \Omega$ whereby $\Omega \neq \sigma^2 \mathbf{I}$ as a known positive definite $n \times n$ matrix
3. $\mathbf{X}$ has full rank.

Since we have that Ω is a positive definite and known matrix, we can find a non-stochastic $n \times n$ matrix $\mathbf{B}$ with properties such that:

$$\mathbf{B}^T \mathbf{B} = \Omega^{-1}$$

and

$$\mathbf{B}\Omega\mathbf{B}^T = \mathbf{I}.$$

The trick is to now transform our model from one that contains heteroskedasticity to one with homoskedasticity.

To this end, let us define the transformed variables:

$$\begin{aligned} \mathbf{y}^* &= \mathbf{B}\mathbf{y} \\ \mathbf{X}^* &= \mathbf{B}\mathbf{X}. \end{aligned}$$

Proposition 242 (Moments of Transformed Regression Model)

The moments of the transformed regression model are:

$$\begin{aligned} \mathbb{E}[\mathbf{y}^*|\mathbf{X}] &= \mathbf{X}^*\boldsymbol{\beta} \\ \text{Var}[\mathbf{y}^*|\mathbf{X}] &= \mathbf{I} \\ \mathbf{X}^* &= \mathbf{B}\mathbf{X} \text{ has full rank.} \end{aligned}$$

Proof. First:

$$\mathbb{E}[\mathbf{y}^*|\mathbf{X}] = \mathbb{E}[\mathbf{B}\mathbf{y}|\mathbf{X}] = \mathbf{B}\mathbb{E}[\mathbf{y}|\mathbf{X}] = \mathbf{B}\mathbf{X}\boldsymbol{\beta} = \mathbf{X}^*\boldsymbol{\beta}.$$

Then:

$$\text{Var}[\mathbf{y}^*|\mathbf{X}] = \text{Var}[\mathbf{B}\mathbf{y}|\mathbf{X}] = \mathbf{B}\text{Var}[\mathbf{y}|\mathbf{X}]\mathbf{B}^T = \mathbf{B}\boldsymbol{\Omega}\mathbf{B}^T = \mathbf{I}$$

by the defining property $\mathbf{B}\boldsymbol{\Omega}\mathbf{B}^T = \mathbf{I}$. □

Remark 243. Since the matrix $\mathbf{B}$ is known and non-singular, conditioning on $\mathbf{X}$ and conditioning on $\mathbf{X}^* = \mathbf{B}\mathbf{X}$ are equivalent.

Hence, we now have a transformed model whereby we have removed the heteroskedasticity and ended up back at homoskedasticity.

Definition 244 (Classical Linear Regression Model with Conditional Homoskedasticity)

The transformed model:

$$\mathbf{y}^* = \mathbf{X}^*\boldsymbol{\beta} + \mathbf{u}^* \quad \text{with } \text{Var}[\mathbf{u}^*|\mathbf{X}^*] = \mathbf{I} \quad (182)$$

is a classical linear regression model with conditional homoskedasticity.

This new transformed model gives us an important result whereby we now have an efficient estimator under this transformation.

Theorem 245 (Aitken's Theorem)

The transformed model:

$$\mathbf{y}^* = \mathbf{X}^*\boldsymbol{\beta} + \mathbf{u}^* \quad \text{with } \text{Var}[\mathbf{u}^*|\mathbf{X}^*] = \mathbf{I} \quad (183)$$

has an OLS estimator $\hat{\boldsymbol{\beta}}_{GLS}$ that is efficient in the class of linear unbiased estimators:

$$\hat{\boldsymbol{\beta}}_{GLS} = (\mathbf{X}^T\boldsymbol{\Omega}^{-1}\mathbf{X})^{-1}\mathbf{X}^T\boldsymbol{\Omega}^{-1}\mathbf{y}. \quad (184)$$

Remark 246. This result says that this estimator has the smallest variance out of all linear unbiased estimators.

Proof. First, recall that $\mathbf{X}^* = \mathbf{B}\mathbf{X}$ and $\mathbf{y}^* = \mathbf{B}\mathbf{y}$. Now, use the formula of OLS:

$$\begin{aligned}\hat{\beta}_{GLS} &= ((\mathbf{X}^*)^T \mathbf{X}^*)^{-1} (\mathbf{X}^*)^T \mathbf{y}^* \\ &= (\mathbf{X}^T \mathbf{B}^T \mathbf{B} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{B}^T \mathbf{B} \mathbf{y} \\ &= (\mathbf{X}^T \Omega^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{B}^T \mathbf{B} \mathbf{y} \\ &= (\mathbf{X}^T \Omega^{-1} \mathbf{X})^{-1} \mathbf{X}^T \Omega^{-1} \mathbf{y}\end{aligned}$$

where we used $\mathbf{B}^T \mathbf{B} = \Omega^{-1}$. □

We seemed to have picked a random matrix $\mathbf{B}$ for our transformation but in fact, we could have picked any matrix proportional to $\mathbf{B}$ and still obtained the same GLS estimator.

Proposition 247 (GLS Estimator Invariant Under Scaling Transformation)

Let $\tilde{\mathbf{B}} = a\mathbf{B}$ for some scalar a . Then, under the transformed variables:

$$\begin{aligned}\tilde{\mathbf{y}} &= \tilde{\mathbf{B}}\mathbf{y} \\ \tilde{\mathbf{X}} &= \tilde{\mathbf{B}}\mathbf{X}\end{aligned}$$

we have the same GLS estimator:

$$\hat{\beta}_{GLS} = (\mathbf{X}^T \Omega^{-1} \mathbf{X})^{-1} \mathbf{X}^T \Omega^{-1} \mathbf{y}. \quad (185)$$

Proof. We simply write out the OLS estimator in the transformed model:

$$\begin{aligned}((\tilde{\mathbf{X}})^T \tilde{\mathbf{X}})^{-1} (\tilde{\mathbf{X}})^T \tilde{\mathbf{y}} &= (a^2 \mathbf{X}^T \mathbf{B}^T \mathbf{B} \mathbf{X})^{-1} a^2 \mathbf{X}^T \mathbf{B}^T \mathbf{B} \mathbf{y} \\ &= (\mathbf{X}^T \Omega^{-1} \mathbf{X})^{-1} \mathbf{X}^T \Omega^{-1} \mathbf{y} = \hat{\beta}_{GLS}\end{aligned}$$

□

Remark 248. This invariance property in fact holds even when we replace $\mathbf{B}$ by $\mathbf{Q}\mathbf{B}$ where $\mathbf{Q}$ is a $n \times n$ matrix with the property that $\mathbf{Q}^T \mathbf{Q} = \mathbf{I}$.

We now impose another assumption.

Assumption 249 (Normality)

We assume normality:

$$\mathbf{y} | \mathbf{X} \sim \mathcal{N}(\mathbf{X}\beta, \Omega). \quad (186)$$

We then get a useful result about the estimator under normality.

Proposition 250 (GLS and MLE)

Under normality of the error term, then the GLS estimator coincides with the conditional maximum likelihood estimator when Ω is known.

Remark 251. This is an analogous result to the fact that OLS coincides with MLE when the error terms are normally distributed.

We can then get a finite sample distribution of the GLS estimator.

Proposition 252 (Finite Sample Normality of GLS)

Under normality of the error term, we have that the finite sample distribution of $\hat{\beta}_{GLS}$ is:

$$\hat{\beta}_{GLS}|\mathbf{X} \sim \mathcal{N}(\beta, (\mathbf{X}^T \Omega^{-1} \mathbf{X})^{-1}) \quad (187)$$

8.2.1 Feasible Generalized Least Squares

In practice, we cannot compute this GLS estimator as we do not know the conditional variance matrix Ω . Therefore, we need to estimate Ω when constructing our estimator.

Definition 253 (Feasible GLS Estimator)

The **feasible generalised least squares** estimator is given by:

$$\hat{\beta}_{FGLS} = (\mathbf{X}^T \hat{\Omega}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \hat{\Omega}^{-1} \mathbf{y} \quad (188)$$

where $\hat{\Omega}$ is an estimator of Ω .

In practice, we can estimate the feasible GLS by again transforming our model and running OLS. The transformation in this case is given by:

$$\begin{aligned} \mathbf{y}^* &= \hat{\mathbf{B}} \mathbf{y} \\ \mathbf{X}^* &= \hat{\mathbf{B}} \mathbf{X} \end{aligned}$$

where we require that $\hat{\mathbf{B}}^T \hat{\mathbf{B}} = \hat{\Omega}^{-1}$ and $\hat{\mathbf{B}} \hat{\Omega} \hat{\mathbf{B}}^T = \mathbf{I}$.

We can derive the properties of the FGLS estimator. However, such properties depend on the properties of $\hat{\Omega}$ as well.

Proposition 254 (Asymptotic Normality of FGLS)

Suppose that $\Omega = \Omega(\phi)$ is correctly specified with a fixed-dimensional parameter ϕ , that $\hat{\phi}$ is consistent, and that the standard regularity conditions hold. Then, $\hat{\beta}_{FGLS}$ has the asymptotic distribution:

$$\hat{\beta}_{FGLS} \overset{a}{\sim} \mathcal{N}\left(\beta, (\mathbf{X}^T \hat{\Omega}^{-1} \mathbf{X})^{-1}\right). \quad (189)$$

Furthermore, $\hat{\beta}_{FGLS}$ is asymptotically efficient.

We can therefore see that consistency of $\hat{\Omega}$ is quite important in order for us to have nice properties for the FGLS. However, obtaining a consistent estimator $\hat{\Omega}$ is not straightforward. The reason is that the $n \times n$ symmetric matrix Ω has $n(n+1)/2$ distinct elements. Even if we restrict our attention that all off-diagonal elements are zero, we will still have n distinct elements in Ω , which we cannot estimate consistently from a sample of size n .

We have just seen the issue for estimating Ω even with many assumptions. Hence, in order for us to actually estimate Ω , we need to specify a parametric model for the conditional variance matrix $\text{Var}[\mathbf{y}|\mathbf{X}] = \Omega$ of the form:

$$\Omega = \Omega(\phi)$$

where the $n \times n$ matrix $\Omega(\phi)$ is a function of the vector ϕ which contains a finite number of parameters that **does not increase** with the sample size. Therefore, we in fact estimate ϕ instead and see if we can consistently estimate ϕ . This gets us around the earlier issue with having too many parameters to estimate that we mentioned earlier.

Therefore, if the specification of the conditional variance matrix $\text{Var}[\mathbf{y}|\mathbf{X}] = \Omega(\phi)$ is correct and we can find a **consistent estimator** $\hat{\phi}$ of the vector ϕ , then we can then use the consistent estimator:

$$\hat{\Omega} = \Omega(\hat{\phi})$$

to obtain the asymptotically efficient FGLS estimator.

8.2.2 Weighted Least Squares

We look at a simple example of computing the FGLS estimator. Let us specify the conditional variance to be:

$$\text{Var}[y_i|\mathbf{X}] = \text{Var}[u_i|\mathbf{X}] = \sigma_i^2$$

to be proportional to the squared values of one of the regressors

$$\sigma_i^2 = \sigma^2 x_{Ki}^2, \quad x_{Ki} \neq 0.$$

We picked the squared value to ensure a non-negative variance, and assume $x_{Ki} \neq 0$ so that the variance is positive and the transformation below is well-defined.

We then define the transformations:

$$y_i^* = \frac{y_i}{x_{Ki}},$$

$$x_{ji}^* = \frac{x_{ji}}{x_{Ki}}$$

for each $j = 1, \dots, K$. Note that under this transformation, the intercept $x_{1i} = 1$, we have that:

$$x_{1i}^* = \frac{x_{1i}}{x_{Ki}} = \frac{1}{x_{Ki}} = (x_{Ki})^{-1}$$

and we also have that:

$$x_{Ki}^* = \frac{x_{Ki}}{x_{Ki}} = 1.$$

Definition 255 (Transformed Model in Weighted Least Squares)

Define the transformations:

$$y_i^* = \frac{y_i}{x_{Ki}},$$

$$x_{ji}^* = \frac{x_{ji}}{x_{Ki}}$$

for each $j = 1, \dots, K$. Then, the transformed model in weighted least squares is:

$$y_i^* = (x_i^*)^T \boldsymbol{\beta} + u_i^*. \quad (190)$$

Remark 256. The intuition to what we are doing here is that we are downweighting the observations where the variance of u_i is relatively high.

We now have a useful result from this!

Proposition 257 (Conditional Homoskedasticity)

Under the weighted least squares transformations, we have conditional homoskedasticity:

$$\text{Var}[y_i^*|\mathbf{X}] = \sigma^2 \quad (191)$$

for all $i = 1, \dots, n$.

Proof. We can write this out:

$$\begin{aligned} \text{Var}[y_i^*|\mathbf{X}] &= \text{Var}[u_i^*|\mathbf{X}] \\ &= \frac{\sigma_i^2}{x_{Ki}^2} \\ &= \frac{\sigma^2 x_{Ki}^2}{x_{Ki}^2} \\ &= \sigma^2. \end{aligned}$$

□

We also give an example with a linear model with one exogenous regressor and an intercept:

$$y_i = \beta_1 + \beta_2 x_{2i} + u_i.$$

Suppose we impose the parametric form:

$$\mathbb{E}[u_i^2|\mathbf{x}_i] = \phi_1 + \phi_2 x_{2i}^2 = \sigma^2(\mathbf{x}_i).$$

This is equivalent to saying:

$$u_i^2 = \phi_1 + \phi_2 x_{2i}^2 + \epsilon_i.$$

We first retrieve the OLS residuals $\hat{u}_i$ from the original regression.

We can retrieve ϕ_1 and ϕ_2 by regressing:

$$\hat{u}_i^2 = \phi_1 + \phi_2 x_{2i}^2 + \epsilon_i.$$

Under correct specification and the usual regularity conditions, the OLS estimators $\hat{\phi}_1$ and $\hat{\phi}_2$ consistently estimate ϕ_1 and ϕ_2 . We then form the fitted conditional variance and standard deviation:

$$\hat{\sigma}_i^2 = \hat{\phi}_1 + \hat{\phi}_2 x_{2i}^2, \quad \hat{\sigma}_i = \sqrt{\hat{\sigma}_i^2},$$

provided the fitted variance is positive.

Here, we weight each observation by a factor which is proportional to $\frac{1}{\hat{\sigma}_i}$. Hence, this is where the name weighted least squares estimator comes from. We give less weight to observations where the variance of u_i is relatively high and more weight to observations where the variance of u_i is relatively low. If our specification for the conditional variance $\text{Var}[\mathbf{y}|\mathbf{X}] = \Omega(\phi)$ is correct, and we estimate $\Omega(\phi)$ consistently, this weighting is the source of the efficiency gain compared to OLS.

However, the feasible GLS estimator is **not** generally the conditional maximum likelihood estimator even under normality of error terms. This is because FGLS uses a consistent estimator of ϕ to construct a consistent estimator

of Ω and then optimises with respect to β only. Meanwhile, the conditional MLE maximises the likelihood function with respect to β and ϕ jointly. These are two different estimators unless we have that:

$$\hat{\phi} = \hat{\phi}_{ML},$$

which is unlikely. As a result, we generally have that:

$$\hat{\beta}_{FGLS} \neq \hat{\beta}_{ML}$$

although they are asymptotically equivalent under correct specification and standard regularity conditions. With a correctly specified conditional mean and well-behaved weights, consistency need not depend on the parametric variance specification being correct, but the efficiency property **does** depend on that specification.

8.3 Clustering

We now look at relaxing the assumption that observations are independent over $i = 1, \dots, n$.

Suppose that each of our n individual observations is allocated to one of $c = 1, \dots, C$ clusters. Let n_c denote the number of observations in cluster c , and let y_{ic} be the dependent variable for observation $i = 1, \dots, n_c$ in cluster c . We define the vector of explanatory variables $\mathbf{x}_{ic}$ and error term u_{ic} similarly.

Definition 258 (Clustered Linear Regression Model)

The clustered linear regression model is given by:

$$y_{ic} = \mathbf{x}_{ic}^T \beta + u_{ic} \quad (192)$$

for $i = 1, \dots, n_c$ and $c = 1, \dots, C$.

We denote n_c as the number of observations in cluster c and we must have that:

$$\sum_{c=1}^C n_c = n.$$

We now define $\mathbf{y}_c$ to be the $n_c \times 1$ vector containing the observations on y_{ic} for the n_c observations in cluster c . We also define the matrix $\mathbf{X}_c$ which is the $n_c \times K$ matrix where each row contains the row vector $\mathbf{x}_{ic}^T$ for an observation in cluster c . Finally, we have the error term vector $\mathbf{u}_c$ as the $n_c \times 1$ vector. We can then define the system of C equations:

$$\mathbf{y}_c = \mathbf{X}_c \beta + \mathbf{u}_c \quad \text{for } c = 1, \dots, C.$$

We can then stack these vectors again to get the $n \times 1$ vectors $\mathbf{y}$ and $\mathbf{u}$ in addition to the $n \times K$ matrix $\mathbf{X}$ to get the model:

$$\mathbf{y} = \mathbf{X} \beta + \mathbf{u}.$$

We can now define the OLS estimator of this model.

Proposition 259 (Clustered OLS)

Let the clustered linear regression model be written as:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}. \quad (193)$$

Then the OLS estimator of $\boldsymbol{\beta}$ is:

$$\hat{\boldsymbol{\beta}}_{OLS} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} = \left(\sum_{c=1}^C \mathbf{x}_c^T \mathbf{x}_c \right)^{-1} \left(\sum_{c=1}^C \mathbf{x}_c^T \mathbf{y}_c \right). \quad (194)$$

The important result from aggregating the variables by clusters is that the observations $(y_{ic}, \mathbf{x}_{ic}^T)$ are now assumed to be **independent across clusters** but not within clusters. We will do something similar later on when we look at panel data. This allows the error terms u_{ic} and u_{jc} to be correlated for different observations in the same cluster c .

Asymptotic results are derived by letting the number of clusters $C \rightarrow \infty$. We state important assumptions that we need to help us derive such results.

Assumption 260 (Strict Exogeneity Assumption)

We assume strict exogeneity of the error terms:

$$\mathbb{E}[\mathbf{u}_c | \mathbf{X}] = \mathbf{0} \quad \text{for } c = 1, \dots, C, \quad (195)$$

$$\mathbb{E}[\mathbf{u}_c \mathbf{u}_d^T | \mathbf{X}] = \mathbf{0} \quad \text{for } c \neq d, \quad (196)$$

while dependence among observations within the same cluster is unrestricted.

We have that under these assumptions:

$$\mathbb{E}[\mathbf{u}\mathbf{u}^T | \mathbf{X}] = \Omega = \begin{bmatrix} \Omega_1 & 0 & \dots & 0 \\ 0 & \Omega_2 & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & \Omega_C \end{bmatrix}$$

where each Ω_c is a symmetric and positive definite $n_c \times n_c$ matrix.

Proposition 261 (Asymptotic Normality)

Suppose clusters are independent, $C \rightarrow \infty$, suitable moment and rank conditions hold, and no cluster dominates the sample. Then we have asymptotic normality and consistency of the OLS estimator:

$$\sqrt{C}(\hat{\boldsymbol{\beta}}_{OLS} - \boldsymbol{\beta}) \xrightarrow{D} \mathcal{N}(0, \mathbf{V}). \quad (197)$$

We can construct a consistent estimator of $\mathbf{V}$ by:

$$\hat{\mathbf{V}} = C(\mathbf{X}^T \mathbf{X})^{-1} \left(\sum_{c=1}^C \mathbf{x}_c^T \hat{\mathbf{u}}_c \hat{\mathbf{u}}_c^T \mathbf{x}_c \right) (\mathbf{X}^T \mathbf{X})^{-1}$$

where $\hat{\mathbf{u}}_c = \mathbf{y}_c - \mathbf{X}_c \hat{\boldsymbol{\beta}}_{OLS}$ is the vector of OLS residuals for the n_c observations in cluster c .

Definition 262 (Cluster-Robust Variance Matrix Estimator)

The variance of $\hat{\beta}_{OLS}$ can be estimated using the **cluster-robust variance matrix** estimator:

$$\frac{\hat{\mathbf{V}}}{C} = (\mathbf{X}^T \mathbf{X})^{-1} \left(\sum_{c=1}^C \mathbf{X}_c^T \hat{u}_c \hat{u}_c^T \mathbf{X}_c \right) (\mathbf{X}^T \mathbf{X})^{-1} \quad (198)$$

This allows us to carry out asymptotically valid inference while allowing for dependence within clusters.

Proposition 263 (Cluster-Robust Variance Matrix Estimator and White's Estimator)

The cluster-robust variance matrix estimator simplifies to White's heteroskedasticity-robust variance matrix estimator in the special case with one observation in each cluster.

9 Endogeneity

9.1 Endogenous Explanatory Variable

In practice, there are many instances when the orthogonality condition $\mathbb{E}[\mathbf{x}_i u_i] = 0$ is violated. This is a big issue as the orthogonality condition was required in order to establish consistency. In particular, if one of the explanatory variables is correlated with the error term u_i , then $\text{Cov}(\mathbf{x}_i, u_i) \neq 0$. In a model with an intercept, so that $\mathbb{E}[u_i] = 0$, this violates both $\mathbb{E}[u_i | \mathbf{x}_i] = 0$ and $\mathbb{E}[\mathbf{x}_i u_i] = 0$, and hence the OLS estimator will generally be both **biased** and **inconsistent**.

There are two situations we will look at where this occurs.

1. One or more of the explanatory variables is measured with error. This is known as **measurement error** or **errors-in-variables**.
2. One or more of the relevant explanatory variables is omitted from the model. This is known as **omitted variable bias**.

9.2 Measurement Error

A common concern in applied econometrics is that relevant explanatory variables may be poorly measured. In particular, **attenuation bias** is a potential cause for concern.

Definition 264 (Attenuation Bias)

Attenuation bias is the biasing of the OLS estimator towards 0.

We can state the condition for when attenuation bias occurs.

Proposition 265 (Condition for Attenuation Bias)

Under classical additive measurement error in a single explanatory variable, the OLS slope estimator is attenuated towards zero asymptotically.

Remark 266. *This asymptotic bias does not disappear in large samples. An exact finite-sample direction requires additional distributional assumptions.*

Fortunately, it is the case that if the dependent variable $\mathbf{y}$ is measured with error, attenuation bias does **not** occur.

Proposition 267 (Mismeasured Dependent Variables does Not Cause Attenuation Bias)

Attenuation bias does **not** occur if the dependent variable $\mathbf{y}$ is measured with error, provided the measurement error in $\mathbf{y}$ is uncorrelated with the correctly measured $\mathbf{X}$.

We now mathematically show this. Assume we have a **single explanatory variable** and no intercept term in our model:

$$y_i^* = x_i^* \beta + u_i \quad \text{for } i = 1, \dots, n \quad (199)$$

whereby y_i^* and x_i^* denote the true values of these variables which we may not observe. To make our analysis easier, we impose the following assumptions.

Assumption 268 (Measurement Error Assumptions)

We assume that the following variables have mean zero:

$$\mathbb{E}[u_i] = 0$$

$$\mathbb{E}[x_i^*] = 0$$

$$\mathbb{E}[y_i^*] = 0$$

for $i = 1, \dots, n$.

Furthermore, we assume that

$$\mathbb{E}[x_i^* u_i] = 0$$

for $i = 1, \dots, n$.

Finally, we have **independent and identically distributed** data.

Remark 269. A motivation for why this may be the case could be that these variables are expressed as deviations from their sample means.

We now show our first result.

Proposition 270 (Mismeasured Dependent Variables does Not Cause Attenuation Bias)

Suppose that there is an additive, mean zero measurement error in the dependent variable only

$$y_i = y_i^* + v_i$$

where y_i is the observed value, y_i^* is the true value, and v_i is the measurement error with $\mathbb{E}[v_i] = 0$ for $i = 1, \dots, n$.

Let the model be:

$$y_i^* = x_i^* \beta + u_i \quad \text{for } i = 1, \dots, n$$

where x_i^* is the true value. Then, there is **no attenuation bias** if the measurement error in y_i is uncorrelated with the correctly measured x_i^* :

$$\mathbb{E}[x_i^* v_i] = 0 \quad \text{for } i = 1, \dots, n.$$

Proof. Start off with the model:

$$y_i^* = x_i^* \beta + u_i$$

and substitute in the mismeasured dependent variable $y_i^* = y_i - v_i$. This results in

$$(y_i - v_i) = x_i^* \beta + u_i$$

or alternatively:

$$y_i = x_i^* \beta + (u_i + v_i)$$

For consistency to hold, we require that x_i^* is uncorrelated with the new error term $(u_i + v_i)$ which contains the true linear model error u_i and the measurement error v_i .

By our measurement error assumptions, we assumed that $\mathbb{E}[x_i^* u_i] = 0$. Furthermore, in the proposition, we also assumed that $\mathbb{E}[x_i^* v_i] = 0$.

Therefore, $\mathbb{E}[x_i^* (u_i + v_i)] = 0$ and therefore the OLS estimator for β is a consistent estimator. $\square$

We now show that there are issues when the explanatory variable is mismeasured.

Proposition 271 (Mismeasured Explanatory Variables Cause Attenuation Bias)

Suppose that there is an additive, mean zero measurement error in the explanatory variable only

$$x_i = x_i^* + e_i$$

where x_i is the observed value, x_i^* is the true value, and e_i is the measurement error with $\mathbb{E}[e_i] = 0$ for $i = 1, \dots, n$. Furthermore, suppose that $\mathbb{E}[x_i^* e_i] = 0$, $\mathbb{E}[u_i e_i] = 0$, and $0 < \text{Var}[e_i] < \infty$.

Let the model be:

$$y_i^* = x_i^* \beta + u_i \quad \text{for } i = 1, \dots, n$$

where y_i^* is the true value. Then, the OLS estimator is **biased towards zero asymptotically**.

Proof. Start off with the model:

$$y_i^* = x_i^* \beta + u_i$$

and substitute in the mismeasured explanatory variable $x_i^* = x_i - e_i$. This results in

$$y_i^* = (x_i - e_i) \beta + u_i$$

or alternatively:

$$y_i^* = x_i \beta + (u_i - e_i \beta)$$

Under the classical measurement-error assumptions, $\text{Cov}(x_i, e_i) = \text{Var}(e_i) > 0$ and

$$\text{Cov}(x_i, u_i - e_i \beta) = -\beta \text{Var}(e_i).$$

This implies a non-zero correlation between x_i and the error term in this model $(u_i - e_i \beta)$ whenever $\beta \neq 0$.

From this, we can show that the OLS estimator of β here is biased. First, we look at the model under a mismeasurement in the explanatory variable:

$$y_i^* = x_i \beta + (u_i - e_i \beta).$$

Suppose that the true value of $\beta > 0$. Then, this implies a **negative correlation** between the observed variable x_i and the error term $(u_i - e_i \beta)$. This causes a **downward bias** in the OLS estimator of β . As $\beta > 0$ was positive, this **downward bias** makes the regression slope flatter and biases the OLS estimator **towards zero**.

Suppose that the true value of $\beta < 0$. Then, this implies a **positive correlation** between the observed variable x_i and the error term $(u_i - e_i \beta)$. This causes a **upward bias** in the OLS estimator of β . As $\beta < 0$ was negative, this **upward bias** makes the regression slope flatter and biases the OLS estimator **towards zero**.

Therefore, under both scenarios, the OLS estimator of β will be **biased towards zero** and this is known as **attenuation bias**. $\square$

We now want to show that a mismeasurement in the explanatory variable also causes the OLS estimator to be inconsistent. To do so, we need some further assumptions.

Assumption 272 (Classical Errors-in-variables Assumption)

For $i = 1, \dots, n$, we assume that:

1. $\mathbb{E}[x_i^* e_i] = 0$ Measurement error is uncorrelated with the true value of x_i^*
2. $\mathbb{E}[u_i e_i] = 0$ Measurement error is uncorrelated with the true model error u_i
3. $\text{Var}[e_i] = \sigma_e^2$ Measurement error has finite variance
4. $\text{Var}[x_i^*] = \sigma_{x^*}^2$ Population variance of the true value x_i^* exists and is finite.

Proposition 273 (Mismeasured Explanatory Variables Causes Inconsistency of OLS)

Assume that the classical errors-in-variables assumption holds. Then, suppose that there is an additive, mean zero measurement error in the explanatory variable only

$$x_i = x_i^* + e_i$$

where x_i is the observed value, x_i^* is the true value, and e_i is the measurement error with $\mathbb{E}[e_i] = 0$ for $i = 1, \dots, n$.

Let the model be:

$$y_i^* = x_i^* \beta + u_i \quad \text{for } i = 1, \dots, n$$

where y_i^* is the true value. Then, the OLS estimator is **inconsistent**.

Proof. Following our earlier derivation when showing the bias of the OLS estimator, we can rewrite the model as:

$$y_i^* = x_i \beta + (u_i - e_i \beta).$$

Then, by the formula of the OLS estimator, we have that:

$$\hat{\beta}_{OLS} = (X^T X)^{-1} X^T y^* = \frac{\sum_{i=1}^n x_i y_i^*}{\sum_{i=1}^n x_i^2} = \frac{\frac{1}{n} \sum_{i=1}^n x_i y_i^*}{\frac{1}{n} \sum_{i=1}^n x_i^2}$$

since x_i and y_i^* are scalars. We now plug in $x_i = x_i^* + e_i$ and $y_i^* = x_i^* \beta + u_i$ into our OLS formula to get:

$$= \frac{\frac{1}{n} \sum_{i=1}^n (x_i^* + e_i)(x_i^* \beta + u_i)}{\frac{1}{n} \sum_{i=1}^n (x_i^* + e_i)^2}$$

We can then expand out the terms:

$$= \frac{\beta \frac{1}{n} \sum_{i=1}^n (x_i^*)^2 + \frac{1}{n} \sum_{i=1}^n x_i^* u_i + \beta \frac{1}{n} \sum_{i=1}^n e_i x_i^* + \frac{1}{n} \sum_{i=1}^n e_i u_i}{\frac{1}{n} \sum_{i=1}^n (x_i^*)^2 + 2 \frac{1}{n} \sum_{i=1}^n (x_i^* e_i) + \frac{1}{n} \sum_{i=1}^n e_i^2}$$

Now, we take limits on both sides and apply the weak law of large numbers:

$$\begin{aligned}
p \lim_{n \rightarrow \infty} \hat{\beta}_{OLS} &= \frac{\beta p \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n (x_i^*)^2 + p \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n x_i^* u_i + \beta p \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n e_i x_i^* + p \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n e_i u_i}{p \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n (x_i^*)^2 + 2p \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n (x_i^* e_i) + p \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n e_i^2} \\
&= \frac{\beta \mathbb{E}[(x_i^*)^2] + \mathbb{E}[x_i^* u_i] + \beta \mathbb{E}[e_i x_i^*] + \mathbb{E}[e_i u_i]}{\mathbb{E}(x_i^*)^2 + 2\mathbb{E}[x_i^* e_i] + \mathbb{E}[e_i^2]} \\
&=^{(1)} \frac{\beta \mathbb{E}[(x_i^*)^2] + 0 + 0 + 0}{\mathbb{E}(x_i^*)^2 + 0 + \mathbb{E}[e_i^2]} \\
&=^{(2)} \left(\frac{\sigma_{x^*}^2}{\sigma_{x^*}^2 + \sigma_e^2} \right) \beta \\
&= \frac{\beta}{1 + (\sigma_e^2 / \sigma_{x^*}^2)} \neq \beta
\end{aligned}$$

if $\sigma_e^2 > 0$ whereby $=_{(1)}$ uses the classical errors-in-variables assumption and $=_{(2)}$ uses the fact that the first moment of the true variable x_i^* is 0.

Therefore, we have the result that:

$$\begin{aligned}
p \lim_{n \rightarrow \infty} \hat{\beta}_{OLS} &= \frac{\beta}{1 + (\sigma_e^2 / \sigma_{x^*}^2)} < \beta \quad \text{for } \beta > 0 \text{ and } \sigma_e^2 > 0 \\
p \lim_{n \rightarrow \infty} \hat{\beta}_{OLS} &= \frac{\beta}{1 + (\sigma_e^2 / \sigma_{x^*}^2)} > \beta \quad \text{for } \beta < 0 \text{ and } \sigma_e^2 > 0
\end{aligned}$$

where we can see that the OLS estimator of β is inconsistent, with a bias towards zero that does not diminish as the sample becomes large. $\square$

For a given level $\sigma_{x^*}^2$, the severity of the **attenuation bias** increases with the variance of the measurement error σ_e^2 . That is, the bias gets worse as there is more variance in the measurement error. Hence, the magnitude of the inconsistency depends inversely on the **signal-to-noise** ratio ($\sigma_{x^*}^2 / \sigma_e^2$).

One thing to be careful of is that under heteroskedastic measurement error, the presence of measurement error may introduce an incorrect indication of non-linearity in the relationship between the true value y_i^* and the observed variable x_i .

9.2.1 Multiple regression with errors in variables

We can now generalise our scenario to a multiple regression with errors in variables. Suppose that we have K explanatory variables now whereby:

$$\begin{aligned}
y_i^* &= (\mathbf{x}_i^*)^T \boldsymbol{\beta} + u_i \\
\mathbf{x}_i^T &= (\mathbf{x}_i^*)^T + \mathbf{e}_i^T \quad (\text{Measurement error in explanatory variables})
\end{aligned}$$

whereby $\mathbf{x}_i^T, (\mathbf{x}_i^*)^T, \mathbf{e}_i^T$ are $1 \times K$ vectors. We still have measurement error in the explanatory variables. We can then substitute the terms in to get:

$$y_i^* = \mathbf{x}_i^T \boldsymbol{\beta} + (u_i - \mathbf{e}_i^T \boldsymbol{\beta})$$

Under classical vector measurement error, we have

$$\mathbb{E}[\mathbf{x}_i (u_i - \mathbf{e}_i^T \boldsymbol{\beta})] = -\mathbb{E}[\mathbf{e}_i \mathbf{e}_i^T] \boldsymbol{\beta},$$

which is generally non-zero. In particular, it is non-zero if the measurement-error covariance matrix is positive definite and $\boldsymbol{\beta} \neq \mathbf{0}$. The OLS estimator of the $K \times 1$ parameter vector $\boldsymbol{\beta}$ will therefore generally be biased and inconsistent.

We have that if **only one** of the explanatory variables x_i is measured with error, we can show that:

1. The OLS estimator of the coefficient on that mismeasured variable is biased towards zero (**attenuation bias**)
2. The OLS estimator of the coefficients on the other explanatory variables are also biased but in **unknown directions**.

If several explanatory variables are measured with error, it is very difficult to sign the biases for any of the coefficients.

9.3 Omitted Variable

Another common concern in applied econometrics is that relevant explanatory variables may be omitted from the model. This can cause the OLS estimator to be **biased** if the omitted variable is relevant and correlated with the included regressor and the OLS is also inconsistent. Finally, the direction of the bias depends on the sign of the correlation between the included variable and the omitted variable and the sign of the coefficient on the omitted variable in the data generating process.

9.3.1 Single Linear Regression with Omitted Variables

First consider a simple model with one included variable x_{1i} and one omitted variable x_{2i} . The true process which generates the outcomes y_i is assumed to be:

$$y_i = x_{1i}\beta_1 + x_{2i}\beta_2 + u_i \quad \text{for } i = 1, \dots, n$$

satisfying

$$\mathbb{E}[u_i] = \mathbb{E}[x_{1i}u_i] = \mathbb{E}[x_{2i}u_i] = 0$$

and

$$\mathbb{E}[x_{1i}u_i] = \mathbb{E}[x_{2i}u_i] = 0$$

Suppose that the model excludes x_{2i} :

$$y_i = x_{1i}\beta_1 + (u_i + x_{2i}\beta_2) \quad \text{for } i = 1, \dots, n$$

We can stack across the n observations to get:

$$\mathbf{y} = \mathbf{X}_1\beta_1 + (\mathbf{u} + \mathbf{X}_2\beta_2)$$

The OLS estimator of β_1 is then given by:

$$\hat{\beta}_1 = (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{y}$$

$$\begin{aligned} \hat{\beta}_1 &= (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{y} \\ &= (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T (\mathbf{X}_1\beta_1 + \mathbf{X}_2\beta_2 + \mathbf{u}) \\ &= \beta_1 + [(\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_2] \beta_2 + (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{u} \\ &= \beta_1 + \delta\beta_2 + (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{u} \end{aligned}$$

Here, we defined:

$$\hat{\delta} \equiv (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_2$$

to be the coefficient δ in a linear projection of the omitted variable x_{2i} on the included variable x_{1i} :

$$x_{2i} = x_{1i}\delta + e_i \quad \text{for } i = 1, \dots, n$$

Here, we have that $\hat{\delta}$ is actually a consistent estimator of δ so we have that:

$$p \lim_{n \rightarrow \infty} \hat{\delta} = \delta.$$

We also use the fact that:

$$\mathbb{E}[\mathbf{X}_1^T \mathbf{u}] = \mathbf{0}$$

to get rid of the third term by the weak law of large numbers.

We go back to our original expression and we take probability limits:

$$p \lim_{n \rightarrow \infty} \hat{\beta}_1 = \beta_1 + (p \lim_{n \rightarrow \infty} \hat{\delta})\beta_2 = \beta_1 + \delta\beta_2$$

Therefore, we have the important expression:

$$p \lim_{n \rightarrow \infty} \hat{\beta}_1 = \beta_1 + \delta\beta_2 \quad (200)$$

We can see two things. The OLS estimator of β_1 in the model that omits x_{2i} is **inconsistent** unless either:

1. $\delta = 0$ The omitted variable is orthogonal to the included variable;
2. $\beta_2 = 0$ The omitted variable is not a relevant explanatory variable in the true model.

We can see the direction of the bias of the OLS estimator for an omitted relevant explanatory variable ($\beta_2 \neq 0$). The sign of the asymptotic bias is the sign of $\delta\beta_2$. Thus, if x_{1i} and x_{2i} are positively correlated, then $\delta > 0$ and the bias is upward when $\beta_2 > 0$ but downward when $\beta_2 < 0$. The signs are reversed when x_{1i} and x_{2i} are negatively correlated, so that $\delta < 0$.

9.3.2 Multiple Regression with Omitted Variables

We now look at the regression case with multiple explanatory variables. Suppose we had K_1 included regressors and K_2 omitted variables. Then, our regression model is:

$$\mathbf{y} = \mathbf{X}_1\boldsymbol{\beta}_1 + (\mathbf{u} + \mathbf{X}_2\boldsymbol{\beta}_2)$$

where $\mathbf{X}_1$ is a $n \times K_1$ matrix, $\boldsymbol{\beta}_1$ is a $K_1 \times 1$ vector, $\mathbf{X}_2$ is a $n \times K_2$ matrix, and $\boldsymbol{\beta}_2$ is a $K_2 \times 1$ vector.

Following the same steps as before, we can obtain the expression for the OLS estimator as:

$$p \lim_{n \rightarrow \infty} \hat{\beta}_1 = \beta_1 + (p \lim_{n \rightarrow \infty} (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_2) \beta_2 \quad (201)$$

whereby

$$\hat{\boldsymbol{\delta}} = (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_2$$

is the $K_1 \times K_2$ estimator of the coefficient matrix in the linear projection:

$$\mathbf{X}_2 = \mathbf{X}_1\boldsymbol{\delta} + \mathbf{e}$$

It is harder to be confident about the direction of the biases now.

However, if there was only one omitted variable that is correlated with several of the included explanatory variables, then the bias in each estimated coefficients will depend on the partial correlation with the omitted variable and not the simple correlation between the included regressor and omitted variable. That is, under the single omitted variable case:

$$y_i = \beta_1 + \beta_2 x_{2i} + \dots + \beta_{K-1} x_{K-1,i} + (\beta_K x_{Ki} + u_i)$$

where $\mathbb{E}[x_{ki} u_i] = 0$ for $k = 1, \dots, K$ and x_{Ki} is the omitted variable, we can show that the source of the bias for each of the OLS estimators is:

$$p \lim_{n \rightarrow \infty} \hat{\beta}_k = \beta_k + (p \lim_{n \rightarrow \infty} \hat{\delta}_k) \beta_K \quad \text{for } k = 1, \dots, K-1$$

whereby $\hat{\delta}_k$ is from the linear projection:

$$X_{Ki} = \delta_1 + \delta_2 x_{2i} + \dots + \delta_{K-1} x_{K-1,i} + v_i$$

whereby $\mathbb{E}[x_{ki} v_i] = 0$ for $k = 1, \dots, K-1$.

If there are several omitted variables, it is hard to predict the direction of the biases. And hence, the OLS estimator will be a biased and inconsistent estimator of the true population parameters.

10 Instrumental Variables

We have seen what are the issues when we have correlation between the error term and at least one of the explanatory variables. Instrumental variables is a method of estimation that we can use to rectify this situation.

Before we proceed, we recap a few things. First, recall that in a **linear projection model**:

$$y = X\beta + u$$

where X is a $N \times K$ matrix (K explanatory variables) we have the **orthogonality condition**:

$$\mathbb{E}[X^T u] = 0$$

where this holds for a system of K equations (for each explanatory variable). Now, if one of our explanatory variables is an intercept term $X_1 = \mathbf{1}_N$, then this orthogonality condition implies:

$$\mathbb{E}[X_1^T u] = \mathbb{E}[\mathbf{1}_N^T u] = 0.$$

Therefore, under i.i.d. sampling, when the regression matrix X contains an intercept, the population projection error has mean zero:

$$\mathbb{E}[u_i] = 0.$$

If we do not have an intercept, however, the error term need not have mean zero:

$$\mathbb{E}[u_i] \text{ need not equal } 0.$$

This was an important result we need when we want to show finite sample properties of OLS. Therefore, if we **do not** include an intercept term, we then need to impose the assumption that:

$$\mathbb{E}[u_i] = 0.$$

So now going back, suppose we have a structural model:

$$y_i = x_i^T \beta + u_i \tag{202}$$

for observations $i = 1, \dots, n$ and x_i is a $K \times 1$ vector for our K explanatory variables. If we had correlation between the explanatory variable x_i and the error term u_i , we would have issues with OLS not being a consistent estimator. We now go on to explain what we can do to solve this.

10.1 Two Stage Least Squares

The general intuition for 2SLS is that we want to construct predicted values of x_i which are not correlated with the error terms. To construct these predicted values, we can project the explanatory variables x_i onto a set of **instrumental variables** which we hope are:

1. **Informative**: Related to the explanatory variable x_i
2. **Valid**: Uncorrelated with the error terms u_i .

If we have these properties, we can replace x_i by these predicted values and therefore guarantee consistency of our OLS estimator $\hat{\beta}$.

We will now go into 2SLS. We make the following assumptions to make life easier by restricting our attention to the case of a single endogenous variable and instrumental variable.

Assumption 274 (Single Explanatory And Instrumental Variable)

We have a single endogenous explanatory variable x_i . We have a single outside instrumental variable z_i . We assume that both have mean zero:

$$\mathbb{E}[x_i] = 0 \quad (203)$$

$$\mathbb{E}[z_i] = 0. \quad (204)$$

We also impose an assumption on the structural error term.

Assumption 275 (Zero Mean Error)

The structural error term has zero mean:

$$\mathbb{E}[u_i] = 0. \quad (205)$$

We shall explain why we have these assumptions later on. First, let us look at our structural model again:

$$y_i = x_i\beta + u_i$$

for $i = 1, \dots, n$. By our assumption, x_i is endogenous which then implies that:

$$\mathbb{E}[x_i u_i] \neq 0.$$

This is an issue as this is the key orthogonality condition we need for the consistency of OLS estimators. We now describe the first stage of the two stage least squares estimator.

Definition 276 (First Stage Linear Projection)

Let x_i be our endogenous variable and z_i be our outside instrumental variable. Define our structural model as

$$y_i = x_i\beta + u_i$$

for $i = 1, \dots, n$. Then, the first stage linear projection is:

$$x_i = z_i\pi + r_i \quad (206)$$

for $i = 1, \dots, n$ where π is a scalar parameter coefficient.

Now, if we look at the first linear projection, we can see where our initial assumptions that $\mathbb{E}[x_i] = \mathbb{E}[z_i] = 0$ came in. If we regressed:

$$x_i = z_i\pi + r_i$$

which **does not** contain an intercept term, we are not guaranteed that the error term r_i has mean zero $\mathbb{E}[r_i] = 0$. Therefore, by assuming both $\mathbb{E}[x_i] = 0$ and $\mathbb{E}[z_i] = 0$, we do not need to include an intercept and we automatically get the result we need for $\mathbb{E}[r_i] = 0$.

We now define the two conditions needed for an instrumental variable: validity and informativeness.

Definition 277 (Valid Instrument)

An instrument z_i is **valid** if it is uncorrelated with the structural error term:

$$\mathbb{E}[z_i u_i] = 0. \quad (207)$$

Definition 278 (Informative Instrument)

An instrument z_i is **informative** if it is correlated with the endogenous variable x_i :

$$\mathbb{E}[z_i x_i] \neq 0. \quad (208)$$

Note that we can rewrite the condition for informativeness in a different manner. Our first stage linear projection was:

$$x_i = z_i \pi + r_i$$

which therefore implies that correlation between z_i and x_i , $\mathbb{E}[z_i x_i] \neq 0$, can also be expressed as:

$$\pi \neq 0.$$

Proposition 279 (Alternative Expression for Informative Instrument)

Suppose an instrument z_i is **informative** for the endogenous variable x_i

$$\mathbb{E}[z_i x_i] \neq 0.$$

if and only if

$$\pi \neq 0 \quad (209)$$

where π is the coefficient in the first stage linear projection:

$$x_i = z_i \pi + r_i$$

for $i = 1, \dots, n$.

Proof. First, recall the first stage linear projection:

$$x_i = z_i \pi + r_i.$$

Now, multiply by z_i on both sides:

$$z_i x_i = z_i^2 \pi + z_i r_i.$$

Now take expectations of both sides:

$$\mathbb{E}[z_i x_i] = \pi \mathbb{E}[z_i^2] + \mathbb{E}[z_i r_i] = \pi \mathbb{E}[z_i^2] + 0$$

where we used the fact that $\mathbb{E}[z_i r_i] = 0$ by definition of a linear projection model. This means that when we run a regression, unless we assume endogeneity, the orthogonality conditions hold. Therefore, we have that:

$$\mathbb{E}[z_i x_i] = \pi \mathbb{E}[z_i^2].$$

By assumption, we said that:

$$\mathbb{E}[z_i x_i] \neq 0$$

as z_i is informative. Now, looking on the right-hand side as $\mathbb{E}[z_i] = 0$ by assumption, we have that:

$$\text{Var}[z_i] = \mathbb{E}[z_i^2].$$

Furthermore, as $\mathbb{E}[z_i x_i] \neq 0$, z_i **must** have a positive variance since a constant cannot be correlated with a random variable $\mathbb{E}[z_i x_i] \neq 0$, therefore we can conclude that:

$$\text{Var}[z_i] = \mathbb{E}[z_i^2] > 0.$$

Therefore, for the expression

$$\mathbb{E}[z_i x_i] = \pi \mathbb{E}[z_i^2]$$

to hold, we therefore require that $\pi \neq 0$ as well since we showed that all the other terms are not equal to zero. $\square$

We can describe the 2SLS estimator. Starting off from our first stage linear projection:

$$x_i = z_i \pi + r_i$$

we can rewrite this in matrix form:

$$\mathbf{X} = \mathbf{Z}\pi + \mathbf{r}$$

where π is still a scalar and $\mathbf{Z}$ is a $n \times 1$ vector. Then, we have that the first stage estimator is given by our usual OLS estimator:

$$\hat{\pi} = (\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T \mathbf{X}.$$

From this, we can now construct our predicted values of the endogenous variable $\mathbf{X} = (x_1, \dots, x_n)^T$ where recall that we only have one explanatory variable. Therefore, if you recall that the predicted value in usual OLS is $\hat{\mathbf{y}} = \mathbf{X}\hat{\boldsymbol{\beta}}$, then the predicted values from our first stage of the endogenous variable is:

$$\hat{\mathbf{X}} = \mathbf{Z}\hat{\pi} = \mathbf{Z} \left[(\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T \mathbf{X} \right] = \mathbf{Z}(\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T \mathbf{X}.$$

We can add and subtract in our predicted value for $\hat{x}_i$ for the endogenous x_i to get the **second stage of the model**:

$$\begin{aligned} y_i &= x_i \beta + u_i \\ \Leftrightarrow y_i &= x_i \beta + u_i + \hat{x}_i \beta - \hat{x}_i \beta \\ \Leftrightarrow y_i &= \hat{x}_i \beta + \left[u_i - \hat{x}_i \beta + x_i \beta \right] \\ \Leftrightarrow y_i &= \hat{x}_i \beta + \left[u_i + (x_i - \hat{x}_i) \beta \right] \end{aligned}$$

Therefore we have the expression:

$$y_i = \hat{x}_i \beta + \left[u_i + (x_i - \hat{x}_i) \beta \right].$$

We can then write this in matrix form:

$$\mathbf{y} = \hat{\mathbf{X}}\boldsymbol{\beta} + (\mathbf{u} + (\mathbf{X} - \hat{\mathbf{X}})\boldsymbol{\beta})$$

where all vectors are $n \times 1$. We can therefore run this regression to get our two-stage least squares (2SLS) estimator.

Definition 280 (First Representation of 2SLS Estimator)

Let x_i be a single endogenous variable and z_i be an instrumental variable. Let our second stage regression be:

$$y = \hat{X}\beta + (u + (X - \hat{X})\beta). \quad (210)$$

Then, the 2SLS estimator is given by:

$$\hat{\beta}_{2SLS} = (\hat{X}^T \hat{X})^{-1} \hat{X}^T y. \quad (211)$$

We can therefore summarise the process for two stage least squares estimation.

Proposition 281 (2SLS Estimation Procedure)

The first stage involves running the regression:

$$X = Z\pi + r \quad (212)$$

and retrieving the first stage estimator:

$$\hat{\pi} = (Z^T Z)^{-1} Z^T X. \quad (213)$$

Then, construct the predicted value:

$$\hat{X} = Z\hat{\pi} = Z(Z^T Z)^{-1} Z^T X. \quad (214)$$

Finally, run the second stage regression:

$$y = \hat{X}\beta + (u + (X - \hat{X})\beta) \quad (215)$$

and retrieve the 2SLS estimator:

$$\hat{\beta}_{2SLS} = (\hat{X}^T \hat{X})^{-1} \hat{X}^T y. \quad (216)$$

Notice that we have another term $(X - \hat{X})\beta$ in the error term. This is in fact a source of finite sample bias.

Proposition 282 (2SLS Estimator is Biased)

The 2SLS estimator is a biased estimator:

$$\mathbb{E}[\hat{\beta}_{2SLS}] \neq \beta. \quad (217)$$

So far, our expression for the 2SLS estimator contains fitted values of $\hat{X}$ and y . We can rewrite our expression for the 2SLS estimator in terms of our original variables X, Z, y .

Proposition 283 (Second Representation of 2SLS Estimator)

The second expression for the 2SLS estimator is:

$$\hat{\beta}_{2SLS} = [X^T Z (Z^T Z)^{-1} Z^T X]^{-1} X^T Z (Z^T Z)^{-1} Z^T y. \quad (218)$$

Proof. First, notice that in our original expression for the 2SLS estimator:

$$\hat{\beta}_{2SLS} = (\hat{X}^T \hat{X})^{-1} \hat{X}^T y$$

we can plug in our expressions for the predicted values:

$$\hat{\mathbf{X}} = \mathbf{Z}\hat{\pi} = \mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X}.$$

We plug this into our original expression to get:

$$\hat{\beta}_{2SLS} = (\hat{\mathbf{X}}^T\hat{\mathbf{X}})^{-1}\hat{\mathbf{X}}^T\mathbf{y} \quad (219)$$

$$= \left(\left[\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X} \right]^T \left[\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X} \right] \right)^{-1} \left[\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X} \right]^T \mathbf{y} \quad (220)$$

$$=_{(1)} \left(\left[\mathbf{X}^T\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T \right] \left[\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X} \right] \right)^{-1} \left[\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X} \right]^T \mathbf{y} \quad (221)$$

$$= \left(\mathbf{X}^T\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X} \right)^{-1} \left[\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X} \right]^T \mathbf{y} \quad (222)$$

$$= \left(\mathbf{X}^T\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{I}_1\mathbf{Z}^T\mathbf{X} \right)^{-1} \left[\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X} \right]^T \mathbf{y} \quad (223)$$

$$= \left(\mathbf{X}^T\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X} \right)^{-1} \left[\mathbf{X}^T\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T \right] \mathbf{y} \quad (224)$$

$$= \left(\mathbf{X}^T\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X} \right)^{-1} \mathbf{X}^T\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{y} \quad (225)$$

$$(226)$$

where in $=_{(1)}$ we applied the transpose and used the symmetry of $((\mathbf{Z}^T\mathbf{Z})^{-1})^T = (\mathbf{Z}^T\mathbf{Z})^{-1}$. $\square$

The previous expression is a bit cumbersome to work with. We can again rewrite our 2SLS estimator.

Proposition 284 (Third Representation of 2SLS Estimator)

The third expression for the 2SLS estimator is:

$$\hat{\beta}_{2SLS} = (\hat{\mathbf{X}}^T\mathbf{X})^{-1}\hat{\mathbf{X}}^T\mathbf{y}. \quad (227)$$

Proof. We take our second expression:

$$\hat{\beta}_{2SLS} = [\mathbf{X}^T\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X}]^{-1}\mathbf{X}^T\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{y}.$$

We then apply the transpose to the whole expression in front of the $\mathbf{y}$:

$$= [(\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X})^T\mathbf{X}]^{-1} \left(\mathbf{Z}(\mathbf{Z}^T\mathbf{Z})^{-1}\mathbf{Z}^T\mathbf{X} \right)^T \mathbf{y}$$

Notice that red and blue terms are exactly the expression for the first representation of the 2SLS estimator:

$$\hat{\beta}_{2SLS} = (\hat{\mathbf{X}}^T\hat{\mathbf{X}})^{-1}\hat{\mathbf{X}}^T\mathbf{y}.$$

We therefore plug this in to get our desired expression:

$$\hat{\beta}_{2SLS} = (\hat{\mathbf{X}}^T\mathbf{X})^{-1}\hat{\mathbf{X}}^T\mathbf{y}.$$

$\square$

Now, we take pause and notice the case for when the 2SLS estimator coincides with the OLS estimator.

Proposition 285 (2SLS and OLS Estimator)

The 2SLS estimator coincides with the OLS estimator:

$$\hat{\beta}_{2SLS} = \hat{\beta}_{OLS} \quad (228)$$

if the instrumental variable used is the explanatory variable itself $z_i = x_i$. This implies that x_i is an exogenous explanatory variable.

Proof. Recall the definition of fitted values for $\hat{x}_i$ and plug in the explanatory variable for the instrumental variable:

$$\hat{\mathbf{X}} = \mathbf{Z}(\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T \mathbf{X} = \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X} = \mathbf{X} \mathbf{I}_1 = \mathbf{X}.$$

Therefore, the predicted value is just the explanatory itself $\hat{\mathbf{X}} = \mathbf{X}$. □

This is very intuitive since if we project x_i on itself, we obtain perfect predictions of x_i and the second stage of 2SLS coincides with the original model.

10.1.1 Just-Identified System

So far, we have assumed only one valid and informative instrument z_i for a **single** endogenous explanatory variable x_i . In such a scenario, we say that the **scalar parameter** β is **just-identified**. For this subsection, we now extend this to K -dimensional parameter vector.

Assumption 286 (K Endogenous Variables and Parameters)

Let β be a $K \times 1$ parameter vector associated to x_i which is a $K \times 1$ vector of explanatory variables.

We now distinguish between the three cases of how many parameters we are estimating and how many instrumental variables we have.

Definition 287 (Just-Identified)

Let z_i contain L valid instruments. The parameter vector β is **just-identified** if $L = K$ and $\text{rank}(\mathbb{E}[z_i x_i^T]) = K$.

Definition 288 (Over-Identified)

The parameter vector β is **over-identified** if $L > K$ and $\text{rank}(\mathbb{E}[z_i x_i^T]) = K$.

Definition 289 (Under-Identified)

The parameter vector β is **under-identified** if $\text{rank}(\mathbb{E}[z_i x_i^T]) < K$. In particular, the order condition $L < K$ implies under-identification.

If we are in the nice case of just-identification, we get a nice expression for the 2SLS estimator.

Proposition 290 (Just-Identified Representation of 2SLS Estimator)

Suppose that the parameter vector β is **just-identified**. Then, the expression for the 2SLS estimator simplifies to:

$$\hat{\beta}_{2SLS} = (\mathbf{Z}^T \mathbf{X})^{-1} \mathbf{Z}^T \mathbf{y}. \quad (229)$$

Proof. We show the proof for our previous setting where $K = 1$. Now, note that $\hat{\pi}$ is a scalar in the first stage linear projection as we only have one endogenous variable: $\hat{\mathbf{X}} = \mathbf{Z}\hat{\pi}$. Then, we rewrite our third representation of the 2SLS estimator using this result:

$$\begin{aligned}\hat{\beta}_{2SLS} &= (\hat{\mathbf{X}}^T \mathbf{X})^{-1} \hat{\mathbf{X}}^T \mathbf{y} \\ &= ((\mathbf{Z}\hat{\pi})^T \mathbf{X})^{-1} (\mathbf{Z}\hat{\pi})^T \mathbf{y} \\ &= (\hat{\pi} \mathbf{Z}^T \mathbf{X})^{-1} \hat{\pi} (\mathbf{Z}^T \mathbf{y}) \\ &=_{(1)} \frac{1}{\hat{\pi}} (\mathbf{Z}^T \mathbf{X})^{-1} \hat{\pi} (\mathbf{Z}^T \mathbf{y}) \\ &= (\mathbf{Z}^T \mathbf{X})^{-1} \mathbf{Z}^T \mathbf{y}\end{aligned}$$

where $=_{(1)}$ makes use of the fact that $\hat{\pi}$ is a scalar. □

10.1.2 L Instruments

So far we looked at the case where we had 1 instrument z_i . We now look at the case where if we had L instruments $\mathbf{z}_i = (z_{1i}, z_{2i}, \dots, z_{Li})^T$ but we still have one endogenous explanatory variable x_i .

Assumption 291 (L Instruments and 1 Endogenous Variable)

Suppose we had L instrumental variables $\mathbf{z}_i = (z_{1i}, z_{2i}, \dots, z_{Li})^T$ and one endogenous explanatory variable x_i .

We again have the structural equation:

$$y_i = x_i \beta + u_i \quad (230)$$

for $i = 1, \dots, n$. To make life easier for us again, we assume that the variables have zero unconditional mean.

Assumption 292 (Zero Unconditional Mean)

Suppose that we had that zero unconditional mean for the structural error term, endogenous variable, and violation of the orthogonality condition:

$$\mathbb{E}[u_i] = \mathbb{E}[x_i] = 0 \quad (231)$$

$$\mathbb{E}[\mathbf{z}_i] = \mathbf{0} \quad (232)$$

$$\mathbb{E}[x_i u_i] \neq 0. \quad (233)$$

Now, since we assumed that $\mathbb{E}[x_i] = 0$ and $\mathbb{E}[\mathbf{z}_i] = \mathbf{0}$, we do not need to include an intercept term in our **first stage linear projection**:

$$x_i = \mathbf{z}_i^T \boldsymbol{\pi} + r_i \quad (234)$$

for $i = 1, \dots, n$. Here, we have 1 endogenous variable x_i but L instruments $\mathbf{z}_i$ and therefore, we have $\boldsymbol{\pi}$ as a $L \times 1$ vector of coefficients on the L instruments. We can then write this in matrix form:

$$\mathbf{X} = \mathbf{Z}\boldsymbol{\pi} + \mathbf{r}. \quad (235)$$

Definition 293 (Validity of L Instruments)

Let $\mathbf{z}_i$ contain L instrumental variables. Then, $\mathbf{z}_i$ is a valid instrument if:

$$\mathbb{E}[\mathbf{z}_i u_i] = \mathbf{0} \quad (236)$$

so that **all** the instrumental variables in $\mathbf{z}_i$ are valid.

Definition 294 (Informativeness of Instruments)

The $\mathbf{z}_i$ instruments are informative for the single explanatory endogenous variable x_i if at least one of the instrumental variables in $\mathbf{z}_i$ is informative:

$$\mathbb{E}[\mathbf{z}_i x_i] \neq \mathbf{0}. \quad (237)$$

If $\mathbb{E}[\mathbf{z}_i \mathbf{z}_i^T]$ is non-singular, equivalently, the coefficient vector on the instruments $\boldsymbol{\pi}$ is not equal to the zero vector:

$$\boldsymbol{\pi} \neq \mathbf{0}. \quad (238)$$

Proof. These two definitions are equivalent for informativeness since first, recall the first stage linear projection:

$$x_i = \mathbf{z}_i^T \boldsymbol{\pi} + r_i.$$

Now, multiply by $\mathbf{z}_i$ on both sides:

$$\mathbf{z}_i x_i = \mathbf{z}_i \mathbf{z}_i^T \boldsymbol{\pi} + \mathbf{z}_i r_i.$$

Taking expectations and using the projection orthogonality condition $\mathbb{E}[\mathbf{z}_i r_i] = \mathbf{0}$ gives:

$$\mathbb{E}[\mathbf{z}_i x_i] = \mathbb{E}[\mathbf{z}_i \mathbf{z}_i^T] \boldsymbol{\pi}.$$

If $\mathbb{E}[\mathbf{z}_i \mathbf{z}_i^T]$ is non-singular, it follows that

$$\mathbb{E}[\mathbf{z}_i x_i] \neq \mathbf{0} \iff \boldsymbol{\pi} \neq \mathbf{0}.$$

□

We can restate the property of informativeness in a different manner.

Proposition 295 (Informativeness Rank Condition)

The $\mathbf{z}_i$ instruments are informative for the single explanatory endogenous variable x_i if:

$$\text{rank}[\boldsymbol{\pi}] = 1 \quad (239)$$

or equivalently,

$$\text{rank}[\mathbb{E}[\mathbf{z}_i x_i]] = 1. \quad (240)$$

We then get two expressions for the 2SLS estimator with L instruments.

Proposition 296 (2SLS Estimator with L Instruments)

Suppose we had L instruments z_i . Then, the 2SLS estimator of β has the form:

$$\hat{\beta}_{2SLS} = (\hat{\mathbf{X}}^T \hat{\mathbf{X}})^{-1} \hat{\mathbf{X}}^T \mathbf{y} \quad (241)$$

and

$$\hat{\beta}_{2SLS} = (\hat{\mathbf{X}}^T \mathbf{X})^{-1} \hat{\mathbf{X}}^T \mathbf{y}. \quad (242)$$

Proof. Plug in the expression that:

$$\hat{\mathbf{X}} = \mathbf{Z}(\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T \mathbf{X}$$

and

$$\hat{\mathbf{X}}^T = \mathbf{X}^T \mathbf{Z}(\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T.$$

Then, follow the steps going from the first representation of the 2SLS estimator to the second representation and then to the third representation. $\square$

10.2 Control Function

We now give another on how we can compute the 2SLS estimator through **control functions**. We make an assumption.

Assumption 297 (L Instruments and 1 Endogenous Variable)

Suppose we had L instrumental variables $\mathbf{z}_i = (z_{1i}, z_{2i}, \dots, z_{Li})^T$ and one endogenous explanatory variable x_i .

We again have the structural equation:

$$y_i = x_i \beta + u_i \quad (243)$$

for $i = 1, \dots, n$ where $\mathbb{E}[u_i] = \mathbb{E}[x_i] = 0$, $\mathbb{E}[\mathbf{z}_i] = \mathbf{0}$, and the orthogonality condition violated $\mathbb{E}[x_i u_i] \neq 0$.

Now, since we assumed that $\mathbb{E}[x_i] = 0$ and $\mathbb{E}[\mathbf{z}_i] = \mathbf{0}$, we do not need to include an intercept term in our **first stage linear projection**:

$$x_i = \mathbf{z}_i^T \boldsymbol{\pi} + r_i \quad (244)$$

for $i = 1, \dots, n$. Here, we have 1 endogenous variable x_i but L instruments $\mathbf{z}_i$ and therefore, we have $\boldsymbol{\pi}$ as a $L \times 1$ vector of coefficients on the L instruments. We can then write this in matrix form:

$$\mathbf{X} = \mathbf{Z} \boldsymbol{\pi} + \mathbf{r}. \quad (245)$$

To make life easier for us again, we assume that the variables have zero unconditional mean.

Assumption 298 (Zero Unconditional Mean)

Suppose that we had that zero unconditional mean for the structural error term, endogenous variable, and violation of the orthogonality condition:

$$\mathbb{E}[u_i] = \mathbb{E}[x_i] = 0 \quad (246)$$

$$\mathbb{E}[\mathbf{z}_i] = \mathbf{0} \quad (247)$$

$$\mathbb{E}[x_i u_i] \neq 0. \quad (248)$$

Now, define the OLS residuals from the first stage linear projection equation as $\hat{r}_i$.

$$\hat{r}_i \equiv x_i - z_i^T \hat{\pi}.$$

Definition 299 (First Stage Linear Projection Residuals)

The OLS residuals from the first stage linear projection equation is defined as:

$$\hat{r}_i \equiv x_i - z_i^T \hat{\pi}. \quad (249)$$

With this, we can define what is the control function approach for 2SLS.

Definition 300 (2SLS Using Control Functions)

Let $\hat{r}_i$ be the OLS residuals from the first stage linear projection equation. Then, the 2SLS estimator $\hat{\beta}_{2SLS}$ can be estimated as the OLS estimator of the extended equation:

$$y_i = x_i \beta + \hat{r}_i \gamma + v_i \quad (250)$$

for $i = 1, \dots, n$. The first stage residuals $\hat{r}_i$ are referred to the **control function**.

First, notice we are now regressing on x_i and not $\hat{x}_i$! The inclusion of the control function allows us to control for correlation between the endogenous x_i and the error term u_i in the original model. The reason it does this is because we can decompose the endogenous variable:

$$x_i = z_i^T \pi + r_i$$

where we assume z_i is uncorrelated with the structural error u_i by definition of an instrumental variable. As z_i is uncorrelated with u_i , then, the only component of x_i that could be correlated with the structural error u_i is therefore r_i . Therefore, by including r_i into our equation, we now control for this correlation between x_i and the structural error u_i . This control function approach gives us a consistent estimator β , which we shall discuss in the next section.

10.3 Indirect Least Squares

For the next method, we need to assume we are in a just-identified setting.

Assumption 301 (Just-Identified)

The parameter vector β is **just-identified** if $L = K$ and $\text{rank}(\mathbb{E}[z_i x_i^T]) = K$.

Suppose that we had the number of parameters is 1, $K = 1$, which implies that we have one endogenous explanatory variable x_i . By just-identified system, this means that the number of instrumental variables is 1, $L = 1$. Now, take our first stage linear projection for x_i :

$$x_i = z_i \pi + r_i$$

and we can then substitute this into our structural equation:

$$y_i = x_i \beta + u_i$$

to get the reduced form equation:

$$y_i = z_i(\pi\beta) + (u_i + r_i\beta).$$

Definition 302 (Reduced Form Equation)

The reduced form equation in indirect least squares is given by:

$$y_i = z_i(\pi\beta) + (u_i + r_i\beta) \quad (251)$$

where number of parameters equal the number of instrumental variables.

In the general case with L instruments and K regressors, the first-stage coefficient matrix Π is $L \times K$ and the reduced-form coefficient $\Pi\beta$ is $L \times 1$. Thus, this multiplication is well defined whether or not $L = K$. Just-identification is required to recover β uniquely by indirect least squares, not for the reduced form to exist.

We can state the indirect least squares method approach.

Definition 303 (Indirect Least Squares)

Let

$$x_i = z_i\pi + r_i \quad (252)$$

be our first stage linear projection equation. Let $\hat{\pi}$ be our OLS estimator in this first stage linear projection equation. Then, let

$$y_i = z_i(\pi\beta) + (u_i + r_i\beta) \quad (253)$$

be the reduced form equation with $\hat{\pi}\beta$ as the OLS estimator in this reduced form equation. Then, the **indirect least squares estimator** $\hat{\beta}_{ID}$ is given by:

$$\hat{\beta}_{ID} = \frac{\hat{\pi}\beta}{\hat{\pi}}. \quad (254)$$

This estimator has some nice properties.

Proposition 304 (Consistency of Indirect Least Squares Estimator)

The indirect least squares estimator $\hat{\beta}_{ID}$ is a consistent estimator.

We can also show that this estimator is equivalent to the 2SLS estimator in a just-identified setting.

Theorem 305 (Indirect Least Squares and 2SLS Estimator)

Suppose that we have a just-identified model. Then, the indirect least squares estimator coincides with the 2SLS estimator.

10.4 Consistency of 2SLS Estimator

First, we recall that in finite samples, the 2SLS estimator is a biased estimator.

Theorem 306 (Bias of 2SLS Estimator)

The 2SLS estimator

$$\hat{\beta}_{2SLS} = (\hat{X}^T X)^{-1} \hat{X}^T y \quad (255)$$

is a **biased estimator** of β .

Whilst we may not have the finite sample property of an unbiased estimator, we at least have a nice asymptotic property for the 2SLS estimator.

Theorem 307 (Consistency of 2SLS Estimator)

Assume that (y_i, x_i, z_i) are independent and identically distributed. Furthermore, suppose that the instrument z_i is:

1. Valid $\mathbb{E}[z_i u_i] = 0$
2. Informative $\mathbb{E}[z_i x_i] \neq 0$ or $\pi \neq 0$.

Then the 2SLS estimator:

$$\hat{\beta}_{2SLS} = (\hat{\mathbf{X}}^T \mathbf{X})^{-1} \hat{\mathbf{X}}^T \mathbf{y} \quad (256)$$

is a **consistent** estimator of β .

Proof. First, write the 2SLS estimator as:

$$\begin{aligned} \hat{\beta}_{2SLS} &= (\hat{\mathbf{X}}^T \mathbf{X})^{-1} \hat{\mathbf{X}}^T \mathbf{y} \\ &= (\hat{\mathbf{X}}^T \mathbf{X})^{-1} \hat{\mathbf{X}}^T (\mathbf{X}\beta + \mathbf{u}) \\ &= (\hat{\mathbf{X}}^T \mathbf{X})^{-1} \hat{\mathbf{X}}^T \mathbf{X}\beta + (\hat{\mathbf{X}}^T \mathbf{X})^{-1} \hat{\mathbf{X}}^T \mathbf{u} \\ &= \beta + (\hat{\mathbf{X}}^T \mathbf{X})^{-1} \hat{\mathbf{X}}^T \mathbf{u}. \end{aligned}$$

We then multiply and divide by n :

$$\begin{aligned} \hat{\beta}_{2SLS} &= \beta + (\hat{\mathbf{X}}^T \mathbf{X})^{-1} \hat{\mathbf{X}}^T \mathbf{u} \\ &= \beta + n(\hat{\mathbf{X}}^T \mathbf{X})^{-1} \left(\frac{1}{n} \right) \hat{\mathbf{X}}^T \mathbf{u} \\ &= \beta + \underbrace{\left(\frac{\hat{\mathbf{X}}^T \mathbf{X}}{n} \right)^{-1}}_{(*)} \underbrace{\left(\frac{\hat{\mathbf{X}}^T \mathbf{u}}{n} \right)}_{(**)} \end{aligned}$$

We first look at (*).

Claim 308.

$$p \lim_{n \rightarrow \infty} \left(\frac{\hat{\mathbf{X}}^T \mathbf{X}}{n} \right)^{-1} = (\pi \mathbb{E}[z_i x_i])^{-1} \neq 0.$$

Proof. We have that (x_i, z_i) is i.i.d. We can then use the expression $\hat{\mathbf{X}} = \mathbf{Z}\hat{\pi}$ and plug this in to get:

$$p \lim_{n \rightarrow \infty} \left(\frac{\hat{\mathbf{X}}^T \mathbf{X}}{n} \right) = p \lim_{n \rightarrow \infty} \left[\hat{\pi} \frac{\mathbf{Z}^T \mathbf{X}}{n} \right] = \pi \mathbb{E}[z_i x_i].$$

Here, $\hat{\pi}$ is a **consistent estimator** of π by the properties of the OLS estimator in the first-stage equation. Furthermore, as z_i is informative for x_i , we have that $\mathbb{E}[z_i x_i] \neq 0$.

We can then apply the **continuous mapping theorem** to conclude that:

$$p \lim_{n \rightarrow \infty} \left(\frac{\hat{\mathbf{X}}^T \mathbf{X}}{n} \right)^{-1} = (\pi \mathbb{E}[z_i x_i])^{-1} \neq 0.$$

□

We now look at (**).

Claim 309.

$$p \lim_{n \rightarrow \infty} \left(\frac{\hat{\mathbf{X}}^T \mathbf{u}}{n} \right) = 0.$$

Proof. We know that $\mathbf{x}_i$ is i.i.d for all $i = 1, \dots, n$. Furthermore, $\mathbb{E}[z_i u_i] = 0$ for all $i = 1, \dots, n$ by validity of the instrumental variable. We can therefore apply the **weak law of large numbers** to get:

$$p \lim_{n \rightarrow \infty} \left(\frac{\hat{\mathbf{X}}^T \mathbf{u}}{n} \right) = p \lim_{n \rightarrow \infty} \left(\hat{\pi} \frac{\mathbf{Z}^T \mathbf{u}}{n} \right) = \pi \mathbb{E}[z_i u_i] =_{(1)} \pi \cdot 0 = 0$$

whereby $=_{(1)}$ arises from applying Slutsky's theorem using both the validity assumption of the instrumental variable $\mathbf{Z}^T$ and consistency of $\hat{\pi} \xrightarrow{P} \pi$. $\square$

Therefore, we can apply the **continuous mapping theorem** to the product of (*) and (**) that we started with, whereby we now have that the **product of the limits is the limit of the products**:

$$p \lim_{n \rightarrow \infty} \left[\underbrace{\beta + \left(\frac{\hat{\mathbf{X}}^T \mathbf{X}}{n} \right)^{-1}}_{(*)} \underbrace{\left(\frac{\hat{\mathbf{X}}^T \mathbf{u}}{n} \right)}_{(**)} \right] = \beta + (\pi \mathbb{E}[z_i x_i])^{-1} \cdot 0 = \beta$$

We can therefore conclude that:

$$p \lim_{n \rightarrow \infty} \hat{\beta}_{2SLS} = \beta$$

Therefore, we have shown consistency. $\square$

Looking at the proof for consistency, we can see that we require our instrumental variables to be both valid and informative in order for our 2SLS estimator $\hat{\beta}_{2SLS}$ to be consistent.

10.5 Asymptotics

We get some nice asymptotic results for 2SLS estimators.

Theorem 310 (Asymptotic Normality of 2SLS Estimators)

Let $\{y_i, x_i, z_i\}$ be i.i.d observations. Suppose that the instruments z_i are both valid and informative. Then,

$$\sqrt{n}(\hat{\beta}_{2SLS} - \beta) \xrightarrow{D} \mathcal{N}(0, V) \quad (257)$$

where V is a variance-covariance matrix.

We can do hypothesis testing using chi-squared Wald tests.

Proposition 311 (Variance Estimator Under Conditional Homoskedasticity)

Under conditional homoskedasticity $\mathbb{E}[u_i^2|z_i] = \sigma^2$ for all $i = 1, 2, \dots, n$, then a consistent estimator of the limit variance matrix V is given by:

$$n\hat{\sigma}^2(\hat{\mathbf{X}}^T \hat{\mathbf{X}})^{-1} \quad (258)$$

where σ^2 can be estimated consistently using the residuals from the 2SLS regression:

$$\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n \hat{u}_i^2 \quad (259)$$

where $\hat{u}_i \equiv y_i - \mathbf{x}_i^T \hat{\beta}_{2SLS}$.

This implies that the variance of the 2SLS estimator is:

$$\hat{Var}[\hat{\beta}_{2SLS}] = \frac{\hat{V}}{n} = \hat{\sigma}^2(\hat{\mathbf{X}}^T \hat{\mathbf{X}})^{-1}. \quad (260)$$

As usual, we can obtain White standard errors for the 2SLS estimator if we are concerned about conditional heteroskedasticity.

Proposition 312 (Variance Estimator Under Conditional Heteroskedasticity)

Under conditional heteroskedasticity $\mathbb{E}[u_i^2|z_i] \neq \sigma^2$ for all $i = 1, 2, \dots, n$, then the variance of $\hat{\beta}_{2SLS}$ is given by:

$$\hat{Var}(\hat{\beta}_{2SLS}) = \frac{\hat{V}}{n} = (\hat{\mathbf{X}}^T \hat{\mathbf{X}})^{-1} \left(\sum_{i=1}^n \hat{u}_i^2 \hat{\mathbf{x}}_i \hat{\mathbf{x}}_i^T \right) (\hat{\mathbf{X}}^T \hat{\mathbf{X}})^{-1} \quad (261)$$

where

$$\hat{u}_i = y_i - \mathbf{x}_i^T \hat{\beta}_{2SLS}. \quad (262)$$

10.6 IV with Multiple Regressors

We now generalise where we have multiple explanatory variables that are endogenous. Suppose we had the parameter vector $\beta = (\beta_1, \beta_2, \dots, \beta_K)$. Let $\mathbf{x}_i^T = (1, x_{2i}, \dots, x_{Ki})$ be a $1 \times K$ row vector of explanatory variables such that:

$$y_i = \mathbf{x}_i^T \beta + u_i \quad (263)$$

where the orthogonality condition no longer holds $\mathbb{E}[\mathbf{x}_i u_i] \neq \mathbf{0}$. Notice that we now have an intercept term in $\mathbf{x}_i$. Furthermore, assume that $\mathbf{z}_i^T$ is a $1 \times L$ row vector containing the intercept, any exogenous explanatory variables contained in $\mathbf{x}_i$, and any additional **excluded instruments** that are valid so that $\mathbb{E}[\mathbf{z}_i u_i] = 0$.

Let M be the number of excluded instruments, let K_{x_1} be the number of included exogenous explanatory variables (including the intercept), and let L be the total number of instruments, so that:

$$L = K_{x_1} + M.$$

We now state the first necessary condition we need for identification.

Definition 313 (Order Condition)

The order condition is that there is at least as many instrumental variables L as there are number of parameters K we are trying to estimate:

$$L \geq K. \quad (264)$$

Recall that instrumental variables here $\mathbf{z}_i$ contains both exogenous explanatory variables and outside instrumental variables.

We now state the necessary and sufficient condition for identification of β .

Definition 314 (Rank Condition)

The **rank condition** is that the $L \times K$ matrix is full rank:

$$\text{rank}(\mathbb{E}[\mathbf{z}_i \mathbf{x}_i^T]) = K. \quad (265)$$

Remark 315. Recall that by linear algebra, $\text{rank}(\mathbb{E}[\mathbf{z}_i \mathbf{x}_i^T]) \leq \min\{K, L\}$. Since $L \geq K$ by the order condition, the rank condition requires this rank to equal K .

The rank condition ensures that the instruments provide K linearly independent sources of variation for the explanatory variables and hence identify β . The absence of perfect multicollinearity among the explanatory variables is a separate requirement.

Let us define the first stage projection equation.

Definition 316 (First Stage Projection)

The first stage projection equation is:

$$\mathbf{X} = \mathbf{Z}\Pi + \mathbf{R} \quad (266)$$

where $\mathbf{X}$ is a $n \times K$ matrix, $\mathbf{Z}$ is a $n \times L$ matrix, Π is a $L \times K$ matrix and $\mathbf{R}$ is a $n \times K$ matrix.

Here, in this first stage projection equation, we have K equations in which we are projecting. We can now state an equivalent condition for the rank condition.

Definition 317 (Equivalent Rank Condition)

If $\mathbb{E}[\mathbf{z}_i \mathbf{z}_i^T]$ is non-singular, an equivalent rank condition is:

$$\text{rank}(\Pi) = K. \quad (267)$$

A few final notes on this section. The parameter vector β is **not identified** if we have one endogenous explanatory variable for which we have no informative excluded instrument. Furthermore, the rank condition also fails if we have one informative instrument for two or more endogenous variables.

We have seen that the rank condition was needed to check for the informativeness of instrumental variables.

10.7 Testing

We are interested in three types of tests for 2SLS estimators. First, some of the instruments are informative. Second, all the instruments valid. Finally, are the explanatory variables exogenous.

10.7.1 Testing for Informativeness of Instrumental Variables

We now want to test whether are our instrumental variables informative for the explanatory variable.

Definition 318 (Testing for Informativeness of Instruments)

Let x_i be a single endogenous variable, let $\mathbf{q}_i$ contain the included exogenous regressors, and let $\mathbf{z}_i$ contain the excluded instruments. Estimate the first-stage linear projection:

$$x_i = \mathbf{q}_i^T \boldsymbol{\delta} + \mathbf{z}_i^T \boldsymbol{\pi} + r_i \quad (268)$$

and run the joint F-test:

$$H_0 : \pi_1 = \pi_2 = \dots = \pi_M = 0.$$

With several endogenous variables, repeat this test for each first-stage equation.

A rule of thumb suggested by Staiger and Stock (1997) and often used in empirical work is that the value of the **F-test statistic should be at least 10** for reliance on the standard asymptotic approximation to the distribution of the 2SLS estimator. For this linear model with a single endogenous explanatory variable and conditional homoskedasticity, Stock and Yogo (2005) provide alternative critical values for this F-test which seek to control either the bias in the 2SLS estimator relative to the bias of the OLS estimator in the same model or the rejection frequency of standard asymptotic t-tests based on the 2SLS estimator.

10.7.2 Test for Validity of Instrumental Variables

The Sargan/Hansen test is a test for validity of instrumental variables.

Assumption 319 (Sargan Test Assumption)

We have a linear model with conditional **homoskedasticity**. Furthermore, we have an over-identified model whereby $L > K$.

Definition 320 (Null Hypothesis for Validity of Instruments)

The null hypothesis for the testing of the validity of instruments $\mathbf{z}_i$ is given by:

$$H_0 : \mathbb{E}[\mathbf{z}_i u_i] = \mathbf{0}. \quad (269)$$

We can now state the Sargan test.

Definition 321 (Sargan Test for Validity)

Under the null hypothesis $H_0 : \mathbb{E}[\mathbf{z}_i u_i] = \mathbf{0}$, the **Sargan test statistic** has the form:

$$J = \left(\frac{1}{\hat{\sigma}^2} \right) \hat{\mathbf{u}}^T \mathbf{Z} (\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T \hat{\mathbf{u}} \xrightarrow{D} \chi^2(L - K) \quad (270)$$

whereby $\hat{u}_i = y_i - \mathbf{x}_i^T \hat{\boldsymbol{\beta}}_{2SLS}$ and $\hat{\sigma}^2$ is a consistent estimator of $\sigma^2 = \mathbb{E}[u_i^2 | \mathbf{z}_i]$.

The intuition behind this test is that the scalar test statistic J is proportional to what we minimise to define the 2SLS estimator via the GMM estimator with $(\mathbf{Z}^T \mathbf{Z})^{-1}$ as the choice for the weight matrix. We try to set the criterion

function of the GMM estimator as close to 0 as possible. A large minimised value of J means that the sample moments cannot be made jointly close to zero, which provides evidence against the overidentifying restrictions.

If we have **conditional heteroskedasticity**, then we should use the **Hansen** test.

10.7.3 Test for Endogeneity

We now want to construct our second test to check for endogeneity of our regressors.

Definition 322 (Null Hypothesis for Endogeneity of Regressors)

The null hypothesis for the testing of the endogeneity of regressors x_i is given by:

$$H_0 : \mathbb{E}[x_i u_i] = \mathbf{0}. \quad (271)$$

Recall that we used instrumental variables because we were worried that our OLS estimators were inconsistent and that by using appropriate instruments, the 2SLS estimator will become consistent. If the instrumental variables were indeed valid and informative, then we would expect the 2SLS estimator to be different to the OLS estimator if one of the explanatory variables were endogenous since the OLS estimator will no longer be consistent. We can therefore compare these two sets of parameter estimates as the basis for the Hausman test.

Under the null hypothesis where we have no endogeneity, then both $\hat{\beta}_{OLS}$ and $\hat{\beta}_{2SLS}$ are consistent estimators. If there was endogeneity, then only $\hat{\beta}_{2SLS}$ will be a consistent estimator. However, under the null, $\hat{\beta}_{2SLS}$ will be less efficient than $\hat{\beta}_{OLS}$.

Definition 323 (Hausman Test for Endogeneity)

Under the null hypothesis $H_0 : \mathbb{E}[x_i u_i] = \mathbf{0}$, the **Hausman test statistic** has the form:

$$h = (\hat{\beta}_{2SLS} - \hat{\beta}_{OLS})^T \left[\widehat{Var}[\hat{\beta}_{2SLS}] - \widehat{Var}[\hat{\beta}_{OLS}] \right]^{-1} (\hat{\beta}_{2SLS} - \hat{\beta}_{OLS}) \xrightarrow{D} \chi^2(K) \quad (272)$$

provided we have **conditional homoskedasticity** $\mathbb{E}[u_i^2 | z_i] = \sigma^2$ and the $K \times K$ variance $\left[\widehat{Var}[\hat{\beta}_{2SLS}] - \widehat{Var}[\hat{\beta}_{OLS}] \right]$ is non-singular.

We can construct an alternative form of the Hausman test.

Definition 324 (Scalar Version of Hausman Test for Endogeneity)

The scalar version of the Hausman test has two forms:

$$h = \frac{(\hat{\beta}_{k,2SLS} - \hat{\beta}_{k,OLS})^2}{[\widehat{Var}[\hat{\beta}_{k,2SLS}] - \widehat{Var}[\hat{\beta}_{k,OLS}]]} \xrightarrow{D} \chi^2(1) \quad (273)$$

and

$$t = \frac{(\hat{\beta}_{k,2SLS} - \hat{\beta}_{k,OLS})}{\sqrt{[\widehat{Var}[\hat{\beta}_{k,2SLS}] - \widehat{Var}[\hat{\beta}_{k,OLS}]]}} \sim \mathcal{N}(0, 1), \quad h = t^2 \quad (274)$$

provided that we have **conditional homoskedasticity** $\mathbb{E}[u_i^2 | z_i] = \sigma^2$.

10.8 Finite Sample Problems

We look at two sources of finite sample bias.

10.8.1 Overfitting

If we use too many instrumental variables z_i relative to the number of endogenous variables x_i , then when we construct our predicted values $\hat{x}_i$, we may end up having a perfect fit so that $\hat{x}_i = x_i$ and therefore we get back our endogenous variable. From this, we did not remove the endogeneity in the system and therefore end up with endogeneity again.

Proposition 325 (Overfitting)

As the number of L instruments approach the sample size n , then the 2SLS estimator $\hat{\beta}_{2SLS}$ tend towards the OLS estimator $\hat{\beta}_{OLS}$.

In such a situation, $\hat{\beta}_{2SLS}$ will have exactly the same bias as $\hat{\beta}_{OLS}$ in the limit as the number of instruments converges to the sample size $L \rightarrow n$. Therefore, if $\hat{\beta}_{2SLS}$ and $\hat{\beta}_{OLS}$ are quite similar, we should check that we are not using too many instruments.

10.8.2 Weak Instrumental Variables

Weak instruments is when the coefficients on the instruments in the first stage equation is barely significant so that the instrument only weakly identifies the endogenous variable x_i . This means that $\pi \approx 0$. In such a situation, we get a finite sample bias in the 2SLS estimator $\hat{\beta}_{2SLS}$ in the direction of the OLS estimator $\hat{\beta}_{OLS}$, which is also biased.

Under conventional fixed-parameter asymptotics, any fixed non-zero relevance is sufficient for consistency, but weak relevance can produce severe finite-sample bias and poor normal approximations. Under local-to-zero weak-instrument asymptotics, the limiting distribution of the 2SLS estimator is non-standard and the estimator can be far away from the true parameter β .

As a result, we should avoid weak instruments whenever possible and should not attach much significance to 2SLS estimates in situations where the instruments are very weak. In these situations, conventional t-tests and Wald tests become quite unreliable.

10.8.3 Anderson-Rubin Test

We now give a test for coefficient hypotheses that remains valid under weak instruments. Suppose that we have the linear model with an intercept x_{1i} and a single endogenous explanatory variable x_{Ki} :

$$y_i = \beta_1 + \beta_2 x_{2i} + \dots + \beta_{K-1} x_{K-1,i} + \beta_K x_{Ki} + u_i \quad (275)$$

with $\mathbb{E}[x_{ki} u_i] = 0$ for all $k = 1, 2, \dots, K - 1$.

Furthermore, suppose that we have M valid excluded instruments $z_{1i}, z_{2i}, \dots, z_{Mi}$ with $\mathbb{E}[z_{mi} u_i] = 0$ for $m = 1, 2, \dots, M$.

We require one of these outside instruments to be informative for the explanatory variable x_{Ki} in the first stage linear projection whereby at least one coefficient on the instruments θ_j is non-zero:

$$x_{Ki} = \delta_1 + \delta_2 x_{2i} + \dots + \delta_{K-1} x_{K-1,i} + \theta_1 z_{1i} + \dots + \theta_M z_{Mi} + r_i. \quad (276)$$

We can run an F-test of $H_0 : \theta_1 = \theta_2 = \dots = \theta_M = 0$ with degrees of freedom $F(M, n - L)$. Failure to reject this null indicates a lack of first-stage relevance and possible under-identification, whereas rejection supports relevance. A commonly used rule of thumb suggested by Staiger and Stock is that the first-stage F-statistic should be at least 10.

We may also be interested in testing coefficients in the presence of weak instruments.

Definition 326 (Anderson Rubin Test)

Suppose we had the structural equation:

$$y_i = \beta_1 x_{1i} + \beta_2 x_{2i} + u_i \quad (277)$$

with one weak instrumental variable z_{1i} available. Suppose we wanted to test $H_0 : \beta_1 = 0$. The Anderson Rubin test says to write the reduced form equation:

$$y_i = (\beta_1 \theta) z_{1i} + (\beta_2 + \beta_1 \delta) x_{2i} + (u_i + \beta_1 r_i) \quad (278)$$

and test the null hypothesis $H_0 : \beta_1 \theta = 0$.

11 Generalised Method of Moments

11.1 Motivating Examples

We give two examples of deriving GMM estimators for linear regression and linear regression with instrumental variables before stating the general framework for GMM.

11.1.1 GMM Applied to Linear Regression

Suppose we have the population model:

$$y_i = \mathbf{x}_i^T \boldsymbol{\beta} + u_i \quad \text{for } i = 1, \dots, n$$

whereby y_i is a scalar, $\mathbf{x}_i$ is a $K \times 1$ vector of explanatory variable and $\boldsymbol{\beta}$ is a $K \times 1$ vector of parameters. Furthermore, it is important to stress that $\boldsymbol{\beta}$ here is the true value of the coefficients of the data generating process for the linear model. We impose the following orthogonality assumption, which we will soon see later is necessary for identification.

Assumption 327 (Orthogonality Condition)

The orthogonality condition is that the explanatory variables are uncorrelated with the true population error terms:

$$\mathbb{E}[\mathbf{x}_i u_i] = 0 \quad \text{for } i = 1, \dots, n$$

Now, the ordinary least squares estimator is the argument to the solution to the following minimisation program:

$$\hat{\boldsymbol{\beta}}_{OLS} \equiv \arg \min_{\tilde{\boldsymbol{\beta}}} \sum_{i=1}^n (y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}})^2$$

where we stress that $\tilde{\boldsymbol{\beta}}$ refers to any particular value that the estimator can take, which is in contrast to the fixed true population parameter value $\boldsymbol{\beta}$.

From our minimisation program, we can solve for the OLS estimator simply by taking first order conditions with respect to $\tilde{\boldsymbol{\beta}}$. We end up with the first order condition:

$$\nabla_{\tilde{\boldsymbol{\beta}}} \sum_{i=1}^n (y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}})^2 = -2 \sum_{i=1}^n \mathbf{x}_i (y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}) = 0$$

where $\nabla_{\tilde{\boldsymbol{\beta}}}$ is the $K \times 1$ gradient vector with partial derivatives with respect to $\tilde{\boldsymbol{\beta}}$.

Divide by 2 on both sides and multiply by $1/n$ gives us:

$$\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i (y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}) = 0$$

What we can then do is to define a function of the **data and parameters** as:

$$g(w_i; \tilde{\boldsymbol{\beta}}) \equiv \mathbf{x}_i (y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}) \quad \text{where } w_i = (y_i, \mathbf{x}_i)$$

where w_i is the observed variables. We use this notation w_i as this can be for any general variables we may observe (later on, w_i will include instrumental variables). We can therefore write our first order condition as the **sample average of a known function of data and parameters**:

$$\frac{1}{n} \sum_{i=1}^n g(w_i, \tilde{\beta}) = 0$$

This sample average in fact as a special name in GMM estimation.

Definition 328 (Sample Moment Condition in Linear Regression)

The **sample moment condition** in linear regression is given by:

$$\frac{1}{n} \sum_{i=1}^n g(w_i, \tilde{\beta}) = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i (y_i - \mathbf{x}_i^T \tilde{\beta}) = 0. \quad (279)$$

Remark 329. The sample moment condition is the equation that defines what our estimator shall be.

As we have a sample moment condition, it makes sense to have a population analogue whereby we replace the sample average with the expectations operator $\mathbb{E}$. This also has an important name too.

Definition 330 (Population Moment Condition in Linear Regression)

The **population moment condition** in linear regression is given by:

$$\mathbb{E}[g(w_i; \tilde{\beta})] = \mathbb{E}[\mathbf{x}_i (y_i - \mathbf{x}_i^T \tilde{\beta})] = 0. \quad (280)$$

We now collect and summarise everything we have done so far before proceeding further. We have assumed that we had a linear model for the data generating process and an orthogonality assumption. With this linear model, we can then construct a population moment condition that needs to hold. We can then construct a corresponding sample moment condition by replacing expectations with averages. Now, to find an estimator of the true population parameter, we can instead find an estimator such that the **sample moment condition holds!**

Proposition 331 (Solution to Population Moment Condition)

The population moment condition is solved:

$$\mathbb{E}[g(w_i; \tilde{\beta})] = \mathbb{E}[\mathbf{x}_i (y_i - \mathbf{x}_i^T \tilde{\beta})] = 0$$

when the estimator $\tilde{\beta} = \beta$.

Proof. Start off with the population moment:

$$\mathbb{E}[g(w_i; \tilde{\beta})] = \mathbb{E}[\mathbf{x}_i (y_i - \mathbf{x}_i^T \tilde{\beta})]$$

Let $\tilde{\beta} = \beta$ from our proposition:

$$\mathbb{E}[\mathbf{x}_i (y_i - \mathbf{x}_i^T \tilde{\beta})] = \mathbb{E}[\mathbf{x}_i (y_i - \mathbf{x}_i^T \beta)]$$

We now want to show that this is equal to zero. Recall that from our linear model specification: $y_i = \mathbf{x}_i^T \beta + u_i$. Plugging this in gives us:

$$\mathbb{E}[\mathbf{x}_i (y_i - \mathbf{x}_i^T \beta)] = \mathbb{E}[\mathbf{x}_i u_i] =_{(1)} 0$$

where $=_{(1)}$ holds by our orthogonality assumption. □

We have shown that $\tilde{\beta} = \beta$ is a **solution** to the population moment condition. We want to know if this solution is also **unique**?

Proposition 332 (Unique Solution to Population Moment Condition)

The estimator $\tilde{\beta} = \beta$ is the **unique** solution to the population moment condition:

$$\mathbb{E}[g(w_i; \tilde{\beta})] = \mathbb{E}[x_i(y_i - x_i^T \tilde{\beta})] = 0$$

if there is no perfect multicollinearity, so that $\text{rank}(\mathbb{E}[x_i x_i^T]) = K$.

Proof. Start off with our population moment condition and plug in our linear model for y_i :

$$\begin{aligned} \mathbb{E}[g(w_i; \tilde{\beta})] &= \mathbb{E}[x_i(y_i - x_i^T \tilde{\beta})] \\ &= \mathbb{E}[x_i(x_i^T \beta + u_i - x_i^T \tilde{\beta})] \\ &= \mathbb{E}[x_i x_i^T \beta + x_i u_i - x_i x_i^T \tilde{\beta}] \\ &= \mathbb{E}[x_i x_i^T (\beta - \tilde{\beta}) + x_i u_i] \\ &= \mathbb{E}[x_i x_i^T (\beta - \tilde{\beta})] + \mathbb{E}[x_i u_i] \\ &=_{(1)} \mathbb{E}[x_i x_i^T (\beta - \tilde{\beta})] + 0 \\ &= \mathbb{E}[x_i x_i^T (\beta - \tilde{\beta})] \end{aligned}$$

where $=_{(1)}$ invokes the orthogonality assumption again. Therefore, we want to know for what values $\tilde{\beta}$ does the following expression hold:

$$\mathbb{E}[x_i x_i^T (\beta - \tilde{\beta})] = 0.$$

We claim that the **only** solution to this is when $\tilde{\beta} = \beta$. Now, we need an important result from linear algebra.

Theorem 333 (Product of Matrix and Vector is Zero)

Suppose that $\mathbf{A}$ is a $M \times N$ matrix. Suppose that $\mathbf{A}$ has $\text{rank}(\mathbf{A}) = L < N$. Then, there are $N - L$ **linearly independent vectors** $\mathbf{b} \in \mathbb{R}^N \setminus \{\mathbf{0}\}$ such that:

$$\mathbf{A}\mathbf{b} = \mathbf{0}.$$

Remark 334. This says that if a matrix $\mathbf{A}$ does not have full **column rank** N and instead had a rank of L , then exists $N - L$ vectors $\mathbf{b}$ such that the product of $\mathbf{A}\mathbf{b}$ is the zero vector $\mathbf{0}$.

This theorem actually has an important corollary which is what we actually need.

Corollary 335

Suppose that $\mathbf{A}$ is a $K \times K$ matrix. Suppose that $\text{rank}(\mathbf{A}) = K$ is full rank. Then, the only vector $\mathbf{b} \in \mathbb{R}^K$ such that:

$$\mathbf{A}\mathbf{b} = \mathbf{0}$$

is when $\mathbf{b} = \mathbf{0}$ is the zero vector.

We now go back to our expression:

$$\mathbb{E}[x_i x_i^T] (\beta - \tilde{\beta})$$

and we can map this to our corollary by letting $\mathbf{A} = \mathbb{E}[\mathbf{x}_i \mathbf{x}_i^T]$ be our $K \times K$ matrix and $\mathbf{b} = \boldsymbol{\beta} - \tilde{\boldsymbol{\beta}}$ be our $K \times 1$ vector. We assumed that $\mathbb{E}[\mathbf{x}_i \mathbf{x}_i^T]$ is full rank. Therefore, by our corollary, the only way for

$$\mathbb{E}[\mathbf{x}_i \mathbf{x}_i^T (\boldsymbol{\beta} - \tilde{\boldsymbol{\beta}})] = 0$$

to hold is if the vector multiplying $\mathbb{E}[\mathbf{x}_i \mathbf{x}_i^T]$ is the zero vector $(\boldsymbol{\beta} - \tilde{\boldsymbol{\beta}}) = \mathbf{0}$ and we only get a zero vector if:

$$\tilde{\boldsymbol{\beta}} = \boldsymbol{\beta}.$$

□

We have just shown the linear regression model can be solved by GMM! The GMM estimator takes our model in terms of its population moment conditions, and solves them by choosing an estimator to set their sample analogue equal to 0.

11.1.2 GMM Applied to Linear Regression with Instrumental Variables

We now derive the GMM framework in the context of linear regression with instrumental variables.

We would like to estimate the parameter $\boldsymbol{\beta}$ in the model:

$$y_i = \mathbf{x}_i^T \boldsymbol{\beta} + u_i \quad \text{for } i = 1, \dots, n. \quad (281)$$

where $\mathbf{x}_i, \boldsymbol{\beta}$ are $K \times 1$ vectors. However, **some** of the explanatory variables $\mathbf{x}_i$ may be **endogenous** $\mathbb{E}[\mathbf{x}_i u_i] \neq \mathbf{0}$. Therefore, we will need to use **instrumental variables**.

Assumption 336 (Instrumental Variable Orthogonality)

An instrumental variable $\mathbf{z}_i$ is one where:

$$\mathbb{E}[\mathbf{z}_i u_i] = \mathbf{0}. \quad (282)$$

Definition 337 (Instrumental Variable)

A variable $\mathbf{z}_i$ is a **valid candidate instrument** if it satisfies the **instrumental variable orthogonality** assumption. A usable set of instruments must additionally satisfy the relevance or rank condition.

Therefore, we have our explanatory variables $\mathbf{x}_i$ and our instrumental variables $\mathbf{z}_i$ whereby:

$$\mathbf{x}_i \in \mathbb{R}^K \quad \text{and} \quad \mathbf{z}_i \in \mathbb{R}^L.$$

First, we note that both $\mathbf{x}_i$ and $\mathbf{z}_i$ will have at least an intercept term common in both variables. Furthermore, any exogenous variables in $\mathbf{x}_i$ will also be included into $\mathbf{z}_i$. Finally, $\mathbf{z}_i$ will have additional outside instrumental variables. We elaborate this more. First, partition the structural equation as:

$$y_i = \mathbf{x}_i^T \boldsymbol{\beta} + u_i = \mathbf{x}_{1i}^T \boldsymbol{\beta}_1 + \mathbf{x}_{2i}^T \boldsymbol{\beta}_2 + u_i$$

whereby we have partitioned the explanatory variables as

$$\mathbf{x}_i = \begin{bmatrix} \mathbf{x}_{1i} \\ \mathbf{x}_{2i} \end{bmatrix} = \begin{array}{l} \text{Exogenous variables, } \mathbb{E}[\mathbf{x}_{1i} u_i] = 0 \\ \text{Endogenous variables, } \mathbb{E}[\mathbf{x}_{2i} u_i] \neq 0 \end{array}$$

Here, $\mathbf{x}_{1i}$ explanatory variables are **valid** instruments for themselves as they are exogenous. We can also partition the instruments $\mathbf{z}_i$ as:

$$\mathbf{z}_i = \begin{bmatrix} \mathbf{z}_{1i} \\ \mathbf{z}_{2i} \end{bmatrix} = \begin{bmatrix} \mathbf{x}_{1i} \\ \mathbf{z}_{2i} \end{bmatrix} = \begin{array}{l} \text{Exogenous variables, } \mathbb{E}[\mathbf{x}_{1i}u_i] = 0 \\ \text{Excluded instruments, } \mathbb{E}[\mathbf{z}_{2i}u_i] = 0 \end{array}$$

The variables $\mathbf{z}_{1i}$ are the exogenous variables $\mathbf{x}_{1i}$ and $\mathbf{z}_{2i} \in \mathbb{R}^M$ are the **excluded** (outside) instruments. These variables are called excluded instruments as they are variables that do not appear in the population model in Eq. (281). To be clear on terminology, $\mathbf{z}_i$ are referred to the **instruments** and $\mathbf{z}_{2i}$ are referred to the **excluded instruments**. Therefore, $\mathbf{z}_i$ contains the intercept term, exogenous explanatory variables and the excluded instruments.

Let us also be clear on the dimensions of all our vectors and matrices. The parameter of interest $\boldsymbol{\beta}$ is a $K \times 1$ vector. The explanatory variables, which contains both exogenous and endogenous variables, is a $K \times 1$ vector. We denote K_{x_1} as the number of exogenous variables (including intercept) and K_{x_2} are the number of endogenous variables such that $K_{x_1} + K_{x_2} = K$. The instrumental variables $\mathbf{z}_i$ is a $L \times 1$ vector which contains the exogenous variables $\mathbf{x}_{1i}$ and the excluded instruments $\mathbf{z}_{2i}$. The excluded instruments $\mathbf{z}_{2i}$ is a $M \times 1$ vector such that $L = M + K_{x_1}$ whereby K_{x_1} is the number of exogenous explanatory variables and L is the total number of instruments.

Now, we wish to derive the population moment conditions to identify the true population parameter $\boldsymbol{\beta}$. We can start off with our instrumental variable orthogonality condition and substitute in our true population model:

$$\begin{aligned} 0 &= \mathbb{E}[\mathbf{z}_i u_i] \\ &= \mathbb{E}[\mathbf{z}_i [y_i - \mathbf{x}_i^T \boldsymbol{\beta}]] \\ &= \mathbb{E}[g(\mathbf{w}_i; \boldsymbol{\beta})] \end{aligned}$$

where we defined $\mathbf{w}_i = (\mathbf{x}_i, y_i, \mathbf{z}_i)$ as our observed variables and $g(\mathbf{w}_i; \tilde{\boldsymbol{\beta}}) = \mathbf{z}_i [y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}]$.

Proposition 338 (Solution to Population Moment Condition)

The population moment condition is solved:

$$\mathbb{E}[g(\mathbf{w}_i; \tilde{\boldsymbol{\beta}})] = \mathbb{E}[\mathbf{z}_i [y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}]] = 0$$

when the estimator $\tilde{\boldsymbol{\beta}} = \boldsymbol{\beta}$.

Proof. Let $\tilde{\boldsymbol{\beta}} = \boldsymbol{\beta}$ and use the IV orthogonality assumption:

$$\begin{aligned} \mathbb{E}[\mathbf{z}_i [y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}]] &= \mathbb{E}[\mathbf{z}_i [y_i - \mathbf{x}_i^T \boldsymbol{\beta}]] \\ &= \mathbb{E}[\mathbf{z}_i u_i] = 0. \end{aligned}$$

□

We can also show that this is an unique solution.

Proposition 339 (Unique Solution to Population Moment Condition)

The estimate $\tilde{\boldsymbol{\beta}} = \boldsymbol{\beta}$ is the **unique solution** to the population moment condition:

$$\mathbb{E}[g(\mathbf{w}_i; \tilde{\boldsymbol{\beta}})] = \mathbb{E}[\mathbf{z}_i [y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}]] = 0$$

if the **rank condition** $\text{rank}(\mathbb{E}[\mathbf{z}_i \mathbf{x}_i^T]) = K$ holds.

Proof. Plug in the true population model and use the IV orthogonality assumption:

$$\begin{aligned}
\mathbb{E}[z_i(y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}})] &= \mathbb{E}[z_i(\mathbf{x}_i^T \boldsymbol{\beta} + u_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}})] \\
&= \mathbb{E}[z_i \mathbf{x}_i^T \boldsymbol{\beta} + z_i u_i - z_i \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}] \\
&= \mathbb{E}[z_i \mathbf{x}_i^T \boldsymbol{\beta} - z_i \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}] + \mathbb{E}[z_i u_i] \\
&= \mathbb{E}[z_i \mathbf{x}_i^T (\boldsymbol{\beta} - \tilde{\boldsymbol{\beta}})] + 0 \\
&= \mathbb{E}[z_i \mathbf{x}_i^T (\boldsymbol{\beta} - \tilde{\boldsymbol{\beta}})]
\end{aligned}$$

Again, we can argue that the only way for $\mathbb{E}[z_i \mathbf{x}_i^T (\boldsymbol{\beta} - \tilde{\boldsymbol{\beta}})] = 0$ is if $\tilde{\boldsymbol{\beta}} = \boldsymbol{\beta}$. The matrix $\mathbb{E}[z_i \mathbf{x}_i^T]$ is $L \times K$, where $L \geq K$, and by assumption it has full column rank K . Therefore, its null space contains only the zero vector, so $\boldsymbol{\beta} - \tilde{\boldsymbol{\beta}} = \mathbf{0}$ and hence $\tilde{\boldsymbol{\beta}} = \boldsymbol{\beta}$. □

Remark 340. The assumption $\text{rank}(\mathbb{E}[z_i \mathbf{x}_i^T]) = K$ was the identification assumption used in the section on instrumental variable estimation.

Now, we want to construct our sample moment conditions from our population moment function $\mathbb{E}[g(w_i; \tilde{\boldsymbol{\beta}})]$.

Definition 341 (Sample Moment Condition in Linear Regression with IV)

The **sample moment condition** in linear regression with instrumental variables is given by:

$$\frac{1}{n} \sum_{i=1}^n g(w_i; \tilde{\boldsymbol{\beta}}) = \frac{1}{n} \sum_{i=1}^n z_i(y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}) = 0. \quad (283)$$

Looking at the sample moment equation, we can see that we have a notable issue. $\mathbf{z}_i \in \mathbb{R}^L$ so then $\mathbf{z}_i(y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}})$ is a $L \times 1$ vector. So we have L equations. However, we only have K parameters we wish to estimate $\boldsymbol{\beta}$. And by assumption, we usually have that $L \geq K$. Therefore, we have an overdetermined system with L equations but K unknowns.

As a result, we have two potential cases that can occur.

Case: $L = K$. Here, we have as many instruments as parameters to identify. This means that the system is **just identified** and we can solve for the K unknowns with the L equations (assuming linearity of the system).

Theorem 342 (Instrumental Variable OLS Estimator)

Assuming a just-identified system where $L = K$, then the instrumental variable OLS estimator is given by:

$$\hat{\boldsymbol{\beta}}_{IV} = (\mathbf{Z}^T \mathbf{X})^{-1} \mathbf{Z}^T \mathbf{y} \quad (284)$$

Proof. First, look at the sample moment condition:

$$\frac{1}{n} \sum_{i=1}^n g(w_i; \tilde{\boldsymbol{\beta}}) = \frac{1}{n} \sum_{i=1}^n z_i(y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}) = 0$$

Drop the n^{-1} scaling:

$$\sum_{i=1}^n z_i(y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}) = \sum_{i=1}^n z_i y_i - \sum_{i=1}^n z_i \mathbf{x}_i^T \tilde{\boldsymbol{\beta}} = 0$$

Re-arrange to get:

$$\sum_{i=1}^n z_i \mathbf{x}_i^T \tilde{\boldsymbol{\beta}} = \sum_{i=1}^n z_i y_i$$

Assuming invertibility of $\sum_{i=1}^n \mathbf{z}_i \mathbf{x}_i^T$, which holds with probability approaching one under the population rank condition, we therefore have that:

$$\hat{\boldsymbol{\beta}}_{IV} = \left(\sum_{i=1}^n \mathbf{z}_i \mathbf{x}_i^T \right)^{-1} \sum_{i=1}^n \mathbf{z}_i y_i$$

or in matrix form:

$$\hat{\boldsymbol{\beta}}_{IV} = (\mathbf{Z}^T \mathbf{X})^{-1} \mathbf{Z}^T \mathbf{y}$$

□

Case: $L > K$. Here, we have more instruments than parameters. This means that the system is **over identified** and we can't solve for the K unknowns with the L equations easily. One way to estimate the parameters is to use the method of **two stage least squares** (2SLS).

Theorem 343 (2SLS Estimator)

Assuming an over identified system where $L > K$, then the Two-Stage Least Squares estimator is given by:

$$\hat{\boldsymbol{\beta}}_{2SLS} = (\hat{\mathbf{X}}^T \mathbf{X})^{-1} \hat{\mathbf{X}}^T \mathbf{y}. \quad (285)$$

Proof. Here, we can take the sample moment conditions:

$$\frac{1}{n} \sum_{i=1}^n g(\mathbf{w}_i; \tilde{\boldsymbol{\beta}}) = \frac{1}{n} \sum_{i=1}^n \mathbf{z}_i (y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}) = 0$$

where we said that $\mathbf{z}_i (y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}})$ was a $L \times 1$ vector. What we can do is to pre-multiply this expression by the first stage regression coefficients $\hat{\boldsymbol{\pi}}_n^T$ and this will reduce the system down to a K dimensional system. We will therefore have K equations and K unknowns and hence we can solve for the estimators explicitly.

First, recall that the first stage regression is given by regressing the explanatory variables on the instrumental variables:

$$\mathbf{X} = \mathbf{Z} \boldsymbol{\Pi} + \mathbf{R}$$

whereby $\boldsymbol{\Pi}$ is a $L \times K$ matrix of the coefficients on the L instruments, $\mathbf{Z}$ is a $n \times L$ matrix containing the observed instruments, and $\mathbf{X}$ is the $n \times K$ matrix of explanatory variables and $\mathbf{R}$ is the error terms. Then, the OLS coefficients in this first stage regression are given by:

$$\hat{\boldsymbol{\Pi}} \equiv (\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T \mathbf{X}$$

whereby $\hat{\boldsymbol{\Pi}}$ is a $L \times K$ matrix. Furthermore, recall that the predicted values of the explanatory variables $\mathbf{X}$ are then:

$$\hat{\mathbf{X}} = \mathbf{Z} \hat{\boldsymbol{\Pi}} = \mathbf{Z} (\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T \mathbf{X}.$$

We can then pre-multiply the sample moment conditions with this first stage coefficients:

$$0 = \frac{1}{n} \hat{\boldsymbol{\Pi}}^T \sum_{i=1}^n \mathbf{z}_i (y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}})$$

We now have a system of K equations and K unknowns! We now recognise that we can use the definition of the predicted values of $\hat{\mathbf{X}}$ to get:

$$= \frac{1}{n} \sum_{i=1}^n \hat{\mathbf{x}}_i (y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}) = \frac{1}{n} \sum_{i=1}^n \hat{\mathbf{x}}_i y_i - \frac{1}{n} \sum_{i=1}^n \hat{\mathbf{x}}_i \mathbf{x}_i^T \tilde{\boldsymbol{\beta}} = 0$$

So we therefore have that:

$$\frac{1}{n} \sum_{i=1}^n \hat{x}_i x_i^T \tilde{\beta} = \frac{1}{n} \sum_{i=1}^n \hat{x}_i y_i$$

We have invertibility of $\sum_{i=1}^n \hat{x}_i x_i^T$ and therefore we get the result:

$$\hat{\beta}_{2SLS} = \left(\frac{1}{n} \sum_{i=1}^n \hat{x}_i x_i^T \right)^{-1} \frac{1}{n} \sum_{i=1}^n \hat{x}_i y_i$$

or in matrix form:

$$\hat{\beta}_{2SLS} = (\hat{X}^T X)^{-1} \hat{X}^T y.$$

□

We can also express the 2SLS estimator in an alternative manner.

Proposition 344 (Alternative Expression for 2SLS)

The 2SLS estimator

$$\hat{\beta}_{2SLS} = (\hat{X}^T X)^{-1} \hat{X}^T y$$

can be expressed as:

$$\hat{\beta}_{2SLS} = (\hat{X}^T \hat{X})^{-1} \hat{X}^T y$$

Proof. Recall that the predicted value is: $\hat{X} = Z(Z^T Z)^{-1} Z^T X$. Then, the transpose is:

$$\hat{X}^T = (Z(Z^T Z)^{-1} Z^T X)^T = X^T Z(Z^T Z)^{-1} Z^T.$$

Now, notice that the two 2SLS estimator expressions differ by $\hat{X}^T X$ and $\hat{X}^T \hat{X}$. We can work backwards to see that:

$$\begin{aligned} \hat{X}^T \hat{X} &= [Z(Z^T Z)^{-1} Z^T X]^T [Z(Z^T Z)^{-1} Z^T X] \\ &= [X^T Z(Z^T Z)^{-1} Z^T] [Z(Z^T Z)^{-1} Z^T X] \\ &= X^T Z(Z^T Z)^{-1} Z^T Z(Z^T Z)^{-1} Z^T X \\ &= X^T Z(Z^T Z)^{-1} I_L Z^T X \\ &= X^T Z(Z^T Z)^{-1} Z^T X \\ &= [Z(Z^T Z)^{-1} Z^T X]^T X \\ &= \hat{X}^T X. \end{aligned}$$

Therefore, we can replace $\hat{X}^T X$ in our original expression of the 2SLS by $\hat{X}^T \hat{X}$ to get the new expression:

$$\hat{\beta}_{2SLS} = (\hat{X}^T \hat{X})^{-1} \hat{X}^T y.$$

□

All of this gives us a nice way to motivate and view that 2SLS as a way of reducing an overdetermined system of sample moment conditions. After doing this reduction, we are then able to explicitly solve this reduced system for the 2SLS estimators.

11.2 GMM General Framework

We can now specify what our general framework is for GMM. Suppose that we have i.i.d. data $w_i \in \mathbb{R}^N$ (this could be $y_i, \mathbf{x}_i, \mathbf{z}_i$). Then, suppose we have a model that is parameterised by the parameter vector $\tilde{\theta} \in \Theta \subseteq \mathbb{R}^K$. Furthermore, there exists a true population parameter vector $\theta \in \Theta$. Notation-wise, $\tilde{\theta}$ refers to any particular value that the population parameter vector can take, whilst θ refers to the unique **true** population parameter implied by the model. The unique true vector θ is what we are trying to estimate. Then, the moment conditions g are a set of functions:

$$g : \mathbb{R}^N \times \Theta \rightarrow \mathbb{R}^L \quad \text{or } (w, \tilde{\theta}) \rightarrow g(w; \tilde{\theta})$$

such that the population moment conditions hold at the true parameter:

$$\mathbb{E}[g(w_i; \theta)] = \mathbf{0}$$

which is a set of L equations. We generally require $L \geq K$ for identification and in the case that $L > K$, then we have an overidentified system whereby we have more moment conditions than parameters. We can then construct sample moment equations based off the population moment conditions:

$$\bar{g}_n(\tilde{\theta}) \equiv \frac{1}{n} \sum_{i=1}^n g(w_i; \tilde{\theta})$$

and we try to set these as close to zero as possible. There are **two ways** to set these sample moment conditions as close as possible to zero in the case of an **overidentified system**.

The first way requires us to define a new type of function.

Definition 345 (Criterion Function)

The criterion function $Q_n(\tilde{\theta})$ is defined to be the scalar:

$$Q_n(\tilde{\theta}) = \bar{g}_n(\tilde{\theta})^T \mathbf{W}_n \bar{g}_n(\tilde{\theta})$$

whereby $\mathbf{W}_n$ is a positive definite $L \times L$ weight matrix and $\bar{g}_n(\tilde{\theta})$ are the sample moment conditions.

We can now define the GMM estimator.

Definition 346 (GMM Estimator)

Given the criterion function Q_n , the GMM estimator $\hat{\theta}_{GMM}$ is defined to be the argument of the minimisation program:

$$\hat{\theta}_{GMM} \equiv \arg \min_{\tilde{\theta} \in \Theta} Q_n(\tilde{\theta})$$

Remark 347. The weight matrix $\mathbf{W}_n$ is how we trade-off between how important it is for 1 moment to be closer to 0 compared to another moment. Choosing different weight matrices $\mathbf{W}_n$ will give us different efficiency properties of the GMM estimator.

The second approach we can choose is to pick a $L \times K$ matrix $\mathbf{L}$ and solve the sample moment conditions:

$$\mathbf{L}^T \bar{g}_n(\tilde{\theta}) = \mathbf{0}$$

whereby we now have a reduced system of K equations. We now have a system of K equations and K unknowns and we can set this system exactly equal to 0.

11.2.1 GMM Applied to IV Model

We take our GMM estimator and apply it to the linear regression model with instrumental variables to see how it works in practice.

Theorem 348 (GMM Estimator with IV Variables)

Under the linear model with instrumental variables, the GMM estimator has the form:

$$\hat{\beta}_{GMM} = \left[(\mathbf{X}^T \mathbf{Z}) \mathbf{W}_n (\mathbf{Z}^T \mathbf{X}) \right]^{-1} (\mathbf{X}^T \mathbf{Z}) \mathbf{W}_n (\mathbf{Z}^T \mathbf{y}) \quad (286)$$

Proof. The population moment conditions are given by:

$$\mathbf{0} = \mathbb{E}[g(\mathbf{w}_i; \boldsymbol{\beta})] = \mathbb{E}[\mathbf{z}_i(y_i - \mathbf{x}_i^T \boldsymbol{\beta})]$$

The sample analogue is given by:

$$\bar{g}_n(\tilde{\boldsymbol{\beta}}) = \frac{1}{n} \sum_{i=1}^n \mathbf{z}_i(y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}) = n^{-1} \mathbf{Z}^T (\mathbf{y} - \mathbf{X} \tilde{\boldsymbol{\beta}})$$

The GMM estimator is defined as:

$$\hat{\beta}_{GMM} = \arg \min_{\tilde{\boldsymbol{\beta}}} Q_n(\tilde{\boldsymbol{\beta}}) = \arg \min_{\tilde{\boldsymbol{\beta}}} \bar{g}_n(\tilde{\boldsymbol{\beta}})^T \mathbf{W}_n \bar{g}_n(\tilde{\boldsymbol{\beta}}).$$

We need this following chain rule for quadratic forms to be used later.

Theorem 349 (Chain Rule of Quadratic Forms)

For any differentiable function $h : \Theta \rightarrow \mathbb{R}^M$,

$$\nabla_{\boldsymbol{\theta}} [h(\boldsymbol{\theta})^T \mathbf{A} h(\boldsymbol{\theta})] = [D_{\boldsymbol{\theta}} h(\boldsymbol{\theta})]^T (\mathbf{A} + \mathbf{A}^T) h(\boldsymbol{\theta}) \quad (287)$$

whereby $\nabla_{\boldsymbol{\theta}}$ is the gradient vector of partial derivatives of $\boldsymbol{\theta}$ and $D_{\boldsymbol{\theta}} h(\boldsymbol{\theta})$ is the Jacobian matrix.

We now compute the first order conditions required for a minimum.

$$\begin{aligned} \mathbf{0} &= \nabla_{\tilde{\boldsymbol{\beta}}} Q_n(\tilde{\boldsymbol{\beta}}) \\ &= \nabla_{\tilde{\boldsymbol{\beta}}} \left[\bar{g}_n(\tilde{\boldsymbol{\beta}})^T \mathbf{W}_n \bar{g}_n(\tilde{\boldsymbol{\beta}}) \right] \\ &=_{(1)} [D_{\tilde{\boldsymbol{\beta}}} \bar{g}_n(\tilde{\boldsymbol{\beta}})]^T (\mathbf{W}_n + \mathbf{W}_n^T) \bar{g}_n(\tilde{\boldsymbol{\beta}}) \\ &=_{(2)} 2 [D_{\tilde{\boldsymbol{\beta}}} \bar{g}_n(\tilde{\boldsymbol{\beta}})]^T \mathbf{W}_n \bar{g}_n(\tilde{\boldsymbol{\beta}}) \\ &=_{(3)} 2 [D_{\tilde{\boldsymbol{\beta}}} (n^{-1} \mathbf{Z}^T (\mathbf{y} - \mathbf{X} \tilde{\boldsymbol{\beta}}))]^T \mathbf{W}_n n^{-1} \mathbf{Z}^T (\mathbf{y} - \mathbf{X} \tilde{\boldsymbol{\beta}}) \\ &= \frac{2}{n} [D_{\tilde{\boldsymbol{\beta}}} (n^{-1} \mathbf{Z}^T \mathbf{y} - n^{-1} \mathbf{Z}^T \mathbf{X} \tilde{\boldsymbol{\beta}})]^T \mathbf{W}_n \mathbf{Z}^T (\mathbf{y} - \mathbf{X} \tilde{\boldsymbol{\beta}}) \\ &= \frac{2}{n} [-n^{-1} \mathbf{Z}^T \mathbf{X}]^T \mathbf{W}_n \mathbf{Z}^T (\mathbf{y} - \mathbf{X} \tilde{\boldsymbol{\beta}}) \\ &= -\frac{2}{n^2} \mathbf{X}^T \mathbf{Z} \mathbf{W}_n \mathbf{Z}^T (\mathbf{y} - \mathbf{X} \tilde{\boldsymbol{\beta}}) \end{aligned}$$

In $=_{(1)}$, we applied the chain rule of quadratic forms. In $=_{(2)}$, we use the fact that $\mathbf{W} = \mathbf{W}^T$ since it is a positive definite matrix. In $=_{(3)}$, we plugged in our definition of sample analogues $\bar{g}_n(\tilde{\boldsymbol{\beta}}) = n^{-1} \mathbf{Z}^T (\mathbf{y} - \mathbf{X}\tilde{\boldsymbol{\beta}})$.

So therefore, we have shown that:

$$\mathbf{0} = -\frac{2}{n^2} \mathbf{X}^T \mathbf{Z} \mathbf{W}_n \mathbf{Z}^T (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}_{GMM}).$$

We then have that:

$$\mathbf{X}^T \mathbf{Z} \mathbf{W}_n \mathbf{Z}^T \mathbf{X} \hat{\boldsymbol{\beta}}_{GMM} = \mathbf{X}^T \mathbf{Z} \mathbf{W}_n \mathbf{Z}^T \mathbf{y}$$

We can then invert the terms on the left hand side to get our final desired result:

$$\hat{\boldsymbol{\beta}}_{GMM} = \left[(\mathbf{X}^T \mathbf{Z}) \mathbf{W}_n (\mathbf{Z}^T \mathbf{X}) \right]^{-1} (\mathbf{X}^T \mathbf{Z}) \mathbf{W}_n (\mathbf{Z}^T \mathbf{y}).$$

□

Proposition 350 (GMM Estimator with the 2SLS Weight Matrix)

Let $\hat{\boldsymbol{\beta}}_{GMM}$ be the GMM estimator in the linear model with instrumental variables. If we set the weight matrix to be:

$$\mathbf{W}_n = [\mathbf{Z}^T \mathbf{Z}]^{-1}$$

then the GMM estimator is equivalent to the Two-Stage Least Squares estimator:

$$\hat{\boldsymbol{\beta}}_{GMM} = [\hat{\mathbf{X}}^T \hat{\mathbf{X}}]^{-1} \hat{\mathbf{X}}^T \mathbf{y} = \hat{\boldsymbol{\beta}}_{2SLS}.$$

Proof. First, notice that if $\mathbf{W}_n = [\mathbf{Z}^T \mathbf{Z}]^{-1}$, then we have that:

$$\mathbf{Z} \mathbf{W}_n \mathbf{Z}^T = \mathbf{Z} [\mathbf{Z}^T \mathbf{Z}]^{-1} \mathbf{Z}^T = \mathbf{P}_Z$$

where $\mathbf{P}_Z$ is the projection matrix onto the subspace spanned by $\mathbf{Z}$. Finally, recall that the predicted values $\hat{\mathbf{X}} = \mathbf{P}_Z \mathbf{X}$.

Therefore, we can plug in our expression for $\mathbf{W}_n = [\mathbf{Z}^T \mathbf{Z}]^{-1}$ and notice the projection matrix $\mathbf{P}_Z$ popping up:

into our expression for the GMM estimator:

$$\begin{aligned} \hat{\boldsymbol{\beta}}_{GMM} &= \left[(\mathbf{X}^T \mathbf{Z}) \mathbf{W}_n (\mathbf{Z}^T \mathbf{X}) \right]^{-1} (\mathbf{X}^T \mathbf{Z}) \mathbf{W}_n (\mathbf{Z}^T \mathbf{y}) \\ &= \left[\mathbf{X}^T (\mathbf{Z} \mathbf{W}_n \mathbf{Z}^T) \mathbf{X} \right]^{-1} \mathbf{X}^T (\mathbf{Z} \mathbf{W}_n \mathbf{Z}^T) \mathbf{y} \\ &= \left[\mathbf{X}^T \mathbf{P}_Z \mathbf{X} \right]^{-1} \mathbf{X}^T \mathbf{P}_Z \mathbf{y} \\ &= \left[\mathbf{X}^T \hat{\mathbf{X}} \right]^{-1} \hat{\mathbf{X}}^T \mathbf{y} = \hat{\boldsymbol{\beta}}_{2SLS}. \end{aligned}$$

□

11.3 Asymptotics

11.3.1 Consistency of GMM Estimator

We now want to show the consistency of the GMM estimator:

$$\hat{\theta}_{GMM} \equiv \arg \min_{\tilde{\theta} \in \Theta} Q_n(\tilde{\theta}) = \arg \min_{\tilde{\theta} \in \Theta} \bar{g}_n(\tilde{\theta})^T \mathbf{W}_n \bar{g}_n(\tilde{\theta})$$

where $\mathbf{W}_n$ is the positive definite weight matrix and

$$\bar{g}_n(\tilde{\theta}) = \frac{1}{n} \sum_{i=1}^n g(w_i; \tilde{\theta})$$

are the sample moments.

We have an issue though. For a GMM estimator, we do not have an explicit formula for the GMM estimator unless we have linear moment conditions. Therefore, it is not so straightforward to derive asymptotic properties for the GMM estimator. However, what we can do is, instead of looking at the GMM estimator, which is defined to be the minimiser of the criterion function $Q_n(\tilde{\theta})$, we can look at the asymptotic properties of the criterion function $Q_n(\tilde{\theta})$ and infer asymptotic properties of the GMM estimator from the asymptotic properties of the criterion function. To proceed, we need two assumptions to help us.

Assumption 351 (GMM-ID)

The true parameter θ uniquely solves the **population moment conditions**:

$$\mathbb{E}[g(w_i; \theta)] = \mathbf{0}.$$

Assumption 352 (GMM-WGHT)

The weight matrix $\mathbf{W}_n$ is a positive definite matrix and

$$\mathbf{W}_n \xrightarrow{P} \mathbf{W}$$

where $\mathbf{W}$ is a positive definite matrix.

We can now state and show our consistency property of the GMM estimator.

Proposition 353 (Consistency of the GMM Estimator)

Under the assumptions [GMM-ID] and [GMM-WGHT], together with standard regularity conditions ensuring uniform convergence of Q_n to a well-separated population criterion, the GMM estimator is a consistent estimator:

$$\hat{\theta}_{GMM} \xrightarrow{P} \theta \tag{288}$$

where θ is the true population parameter vector.

Proof. (Heuristic Proof). Unfortunately, we cannot do a proper proof since we need to actually show a form of uniform convergence, but we sketch a heuristic proof.

First, for each fixed parameter $\tilde{\theta} \in \Theta$, we have that by the **weak law of large numbers**, the sample moments converge to the population moments:

$$\bar{g}_n(\tilde{\theta}) = \frac{1}{n} \sum_{i=1}^n g(w_i; \tilde{\theta}) \xrightarrow{P} \mathbb{E}[g(w_i; \tilde{\theta})]$$

We define:

$$g_0(\tilde{\theta}) \equiv \mathbb{E}[g(w_i; \tilde{\theta})]$$

to be the population moments. We note that by [GMM-ID] assumption, $g_0(\tilde{\theta}) \neq \mathbf{0}$ if $\tilde{\theta} \neq \theta$.

By the **continuous mapping theorem**, for each parameter $\tilde{\theta} \in \Theta$, we have that the sample criterion function converges to the population criterion function:

$$Q_n(\tilde{\theta}) = \bar{g}_n^T(\tilde{\theta}) \mathbf{W}_n \bar{g}_n(\tilde{\theta}) \xrightarrow{P} g_0^T(\tilde{\theta}) \mathbf{W} g_0(\tilde{\theta}) =: Q(\tilde{\theta})$$

where we invoked [GMM-WGHT] assumption for the convergence of $\mathbf{W}$. By [GMM-ID] assumption and the fact that the weight matrix $\mathbf{W}$ is positive definite, we have that:

$$Q(\tilde{\theta}) = \begin{cases} 0 & \text{if } \tilde{\theta} = \theta \\ > 0 & \text{if } \tilde{\theta} \neq \theta \end{cases}$$

We can see that the population criterion function $Q(\tilde{\theta})$ will therefore have a minimum at the true population parameter θ since at θ , we have that $Q(\theta) = 0$ and otherwise, $Q(\tilde{\theta}) > 0$ for $\tilde{\theta} \neq \theta$.

We have shown that the sample criterion function Q_n converges to the population criterion function Q . We heuristically argue that this implies that the GMM estimator $\hat{\theta}_{GMM}$ which is the minimiser of the sample criterion function Q_n (pointwise) converges to the true population parameter θ which is the minimiser of the population criterion function Q . Hence, we have shown our result that:

$$\hat{\theta}_{GMM} \xrightarrow{P} \theta$$

□

11.3.2 Asymptotic Normality of GMM Estimator

We want to be able to do inference on the models we are estimating and hence we need to be able to say something about the distribution of the estimators. We require two regularity conditions.

Assumption 354 (GMM-DIFF)

The moment conditions $\tilde{\theta} \rightarrow g(w; \tilde{\theta})$ are continuously differentiable for all $w \in \mathbb{R}^N$.

Assumption 355 (GMM-INTR)

The true population parameter is in the interior of the parameter space $\theta \in \text{int}(\Theta)$.

[GMM-DIFF] will allow us to compute the first order conditions needed when computing a GMM estimator. [GMM-INTR] will ensure that our estimator will lie in the interior of the parameter space and hence ensure we have an interior solution for the first order conditions. This is because we may have some issues if our minimum is on the boundary of the parameter space as then the first order conditions may not necessarily hold.

We state two more assumptions we will invoke later and explain why we need them later.

Assumption 356 (GMM-JAC)

The Jacobian of the population moments under true values is full rank:

$$\text{Rank}(\mathbf{D}) = \text{Rank}(\mathbf{D}_\theta \mathbb{E}[g(w_i; \theta)]) = K.$$

Assumption 357 (GMM-VAR)

The population covariance matrix $\mathbf{S}$ of the moment equations $g(w_i; \theta)$ under the true parameter θ :

$$\mathbf{S} \equiv \mathbb{E}[g(w_i; \theta)g(w_i; \theta)^T] \quad (289)$$

is a positive definite matrix.

We can now finally state the limiting distribution of the GMM estimator.

Theorem 358 (Asymptotic Normality of GMM Estimator)

Under the assumptions GMM-INTR, GMM-DIFF, GMM-ID, GMM-JAC, GMM-VAR, GMM-WGHT, the GMM estimator $\hat{\theta}_{GMM}$ is **asymptotically normal**:

$$n^{1/2}[\hat{\theta}_{GMM}(\mathbf{W}_n) - \theta] \xrightarrow{D} \mathcal{N}(\mathbf{0}, \mathbf{V}_W) \quad (290)$$

where the variance matrix is given by:

$$\mathbf{V}_W = (\mathbf{D}^T \mathbf{W} \mathbf{D})^{-1} \mathbf{D}^T \mathbf{W} \mathbf{S} \mathbf{W} \mathbf{D} (\mathbf{D}^T \mathbf{W} \mathbf{D})^{-1}.$$

Proof. First, recall that the GMM estimator satisfies the first order conditions when minimising the criterion function $Q_n(\tilde{\theta})$:

$$\mathbf{0} = \nabla_{\tilde{\theta}} Q_n(\tilde{\theta}) = \nabla_{\tilde{\theta}} \left[\bar{g}_n(\tilde{\theta})^T \mathbf{W}_n \bar{g}_n(\tilde{\theta}) \right] = 2[\mathbf{D}_{\tilde{\theta}} \bar{g}_n(\tilde{\theta})]^T \mathbf{W}_n \bar{g}_n(\tilde{\theta}).$$

The GMM estimator $\hat{\theta}_{GMM}$ satisfies the aforementioned first order condition:

$$\mathbf{0} = 2[\mathbf{D}_{\theta} \bar{g}_n(\hat{\theta}_{GMM})]^T \mathbf{W}_n \bar{g}_n(\hat{\theta}_{GMM}) \quad (291)$$

Clearly, we can drop the constant 2 on both sides of the equation. We wish to derive the asymptotic distribution of the GMM estimator from Eq. (291). However, an issue is that Eq. (291) is a **nonlinear** equation and this makes things tricky. Therefore, as we will soon see, we will apply a linearisation to Eq. (291) in order to help us figure out the asymptotic distribution of the GMM estimator. First, let us write Eq. (291) as the product of three terms, which we will analyse the asymptotic properties of separately.

$$[\mathbf{D}_{\theta} \bar{g}_n(\hat{\theta}_{GMM})]^T \mathbf{W}_n \bar{g}_n(\hat{\theta}_{GMM}) = \underbrace{[\mathbf{D}_{\theta} \bar{g}_n(\hat{\theta}_{GMM})]^T}_{(A)} \times \underbrace{\mathbf{W}_n}_{(B)} \times \underbrace{\bar{g}_n(\hat{\theta}_{GMM})}_{(C)}$$

We now analyse each of the three terms separately.

Claim 359. *The expression (A) converges to the Jacobian of the population moment conditions:*

$$(A) = \mathbf{D}_{\theta} \bar{g}_n(\hat{\theta}_{GMM}) \xrightarrow{P} \mathbf{D}_{\theta} \mathbb{E}[g(w_i; \theta)]$$

or equivalently:

$$\mathbf{D}_{\theta} \bar{g}_n(\hat{\theta}_{GMM}) = \mathbf{D}_{\theta} \mathbb{E}[g(w_i; \theta)] + o_p(1)$$

Proof. First, we explain what the expression (A) is:

$$\mathbf{D}_{\theta} \bar{g}_n(\hat{\theta}_{GMM})$$

This is the Jacobian of the **sample moments** evaluated at the **GMM estimator**. Before analysing its asymptotics, we look at the similar Jacobian of the **sample moments** evaluated at a general parameter $\tilde{\theta} \in \Theta$:

$$\begin{aligned} D_{\theta} \bar{g}_n(\tilde{\theta}) &= D_{\theta} \frac{1}{n} \sum_{i=1}^n g(w_i; \tilde{\theta}) \\ &= \frac{1}{n} \sum_{i=1}^n D_{\theta} g(w_i; \tilde{\theta}) \end{aligned}$$

which holds for a fixed $\tilde{\theta} \in \Theta$. Now, we noticed that as $\tilde{\theta}$ is fixed, and w_i is i.i.d, this implies that $g(w_i; \tilde{\theta})$ is a sequence of i.i.d random variables. Therefore, we can apply the **weak law of large numbers**:

$$\frac{1}{n} \sum_{i=1}^n D_{\theta} g(w_i; \tilde{\theta}) \xrightarrow{P} \mathbb{E}[D_{\theta} g(w_i; \tilde{\theta})] \stackrel{(1)}{=} D_{\theta} \mathbb{E}[g(w_i; \tilde{\theta})] = D_{\theta} g_0(\tilde{\theta})$$

whereby we interchanged the expectation and differentiation operator by the **dominating convergence theorem**. This holds for any $\tilde{\theta} \in \Theta$ but we were interested for whether does it holds for our GMM estimator. We know that $\hat{\theta}_{GMM} \xrightarrow{P} \theta$ is a consistent estimator so we can plug in the GMM estimator our the expression we just derived and therefore conclude that:

$$D_{\theta} \bar{g}_n(\hat{\theta}_{GMM}) \xrightarrow{P} D_{\theta} g_0(\theta) \equiv D$$

which says that the Jacobian of the sample moment conditions with the GMM estimator converges to the Jacobian of the population moment conditions at the true parameter, which we denote by D . Equivalently, by little oh-pee notation, we could have written this as:

$$D_{\theta} \bar{g}_n(\hat{\theta}_{GMM}) = D + o_p(1).$$

□

Claim 360. *The expression*

$$(B) = W_n \xrightarrow{P} W.$$

Proof. By the assumption [GMM-WGHT], we immediately have that:

$$W_n \xrightarrow{P} W.$$

□

Claim 361. *The sample moments evaluated at the GMM estimator has a first order Taylor expansion around the true population parameter θ given by:*

$$(C) = \bar{g}_n(\hat{\theta}_{GMM}) = \bar{g}_n(\theta) + (D_{\theta} g_0(\theta) + o_p(1))(\hat{\theta}_{GMM} - \theta)$$

Proof. First, we require a stochastic first-order Taylor expansion theorem.

Proposition 362 (Stochastic First-Order Taylor Expansion)

Suppose that $\theta \rightarrow h(w; \theta)$ is twice continuously differentiable for every $w \in \mathbb{R}^N$ and that $\sup_{\theta \in \Theta} \mathbb{E}[|D_{\theta} h(w_i; \theta)|] < \infty$. Then, for every $\tilde{\theta}_n \xrightarrow{P} \theta \in \text{int}(\Theta)$, we have that:

$$\bar{h}_n(\tilde{\theta}_n) - \bar{h}_n(\theta) = [D_{\theta} \bar{h}_n(\theta) + o_p(1)](\tilde{\theta}_n - \theta) \tag{292}$$

whereby $\bar{h}_n(\theta) = \frac{1}{n} \sum_{i=1}^n h(w_i; \theta)$.

We can immediately apply this Taylor expansion to $\bar{g}_n(\hat{\theta}_{GMM})$ around the true population parameter θ to get:

$$\bar{g}_n(\hat{\theta}_{GMM}) - \bar{g}_n(\theta) = [D_\theta \bar{g}_n(\theta) + o_p(1)](\hat{\theta}_{GMM} - \theta)$$

However, we have already shown that:

$$D_\theta \bar{g}_n(\theta) \xrightarrow{P} D_\theta g_0(\theta) \equiv D$$

and hence we have the result:

$$\bar{g}_n(\hat{\theta}_{GMM}) - \bar{g}_n(\theta) = [D + o_p(1)](\hat{\theta}_{GMM} - \theta)$$

□

We collect our three expressions and hence we have that our first order condition:

$$0 = [D_\theta \bar{g}_n(\hat{\theta}_{GMM})]^T \mathbf{W}_n \bar{g}_n(\hat{\theta}_{GMM}) = \underbrace{[D_\theta \bar{g}_n(\hat{\theta}_{GMM})]^T}_{(A)} \times \underbrace{\mathbf{W}_n}_{(B)} \times \underbrace{\bar{g}_n(\hat{\theta}_{GMM})}_{(C)}$$

can be written as:

$$0 = \left[D + o_p(1) \right]^T \mathbf{W}_n \left[\bar{g}_n(\theta) + [D + o_p(1)](\hat{\theta}_{GMM} - \theta) \right]$$

Now, we restrict our attention to the sample moment averages at the true parameter:

$$\bar{g}_n(\theta) = \frac{1}{n} \sum_{i=1}^n g(w_i; \theta)$$

Clearly, by definition, the moments have population mean 0

$$\mathbb{E}[g(w_i; \theta)] = 0$$

by the GMM-ID assumption. Therefore, we have a sample average with mean 0. We can divide by $\sqrt{n}$ and invoke the **central limit theorem**

$$n^{1/2} \bar{g}_n(\theta) = \frac{1}{n^{1/2}} \sum_{i=1}^n g(w_i; \theta) \xrightarrow{D} \mathcal{N}(0, S)$$

whereby $S \equiv \mathbb{E}[g(w_i; \theta)g(w_i; \theta)^T]$ is the variance-covariance matrix of the sample moments as the expectation is 0. Therefore, we take our original expression and multiply by $n^{1/2}$ and re-arrange to get:

$$\begin{aligned} n^{1/2}(\hat{\theta}_{GMM} - \theta) &= - \left[[D + o_p(1)]^T \mathbf{W}_n [D + o_p(1)] \right]^{-1} \times [D + o_p(1)]^T \mathbf{W}_n n^{1/2} \bar{g}_n(\theta) \\ &= - \left[D^T \mathbf{W}_n D + D^T \mathbf{W}_n o_p(1) + o_p(1) \mathbf{W}_n D + o_p(1) \mathbf{W}_n o_p(1) \right]^{-1} \times \left(D^T \mathbf{W}_n n^{1/2} \bar{g}_n(\theta) + o_p(1) \mathbf{W}_n n^{1/2} \bar{g}_n(\theta) \right) \\ &= - [D^T \mathbf{W}_n D]^{-1} \times D^T \mathbf{W}_n n^{1/2} \bar{g}_n(\theta) + o_p(1) \end{aligned}$$

So we have that:

$$n^{1/2}(\hat{\theta}_{GMM} - \theta) = - [D^T \mathbf{W}_n D]^{-1} \times D^T \mathbf{W}_n n^{1/2} \bar{g}_n(\theta) + o_p(1)$$

By **Slutsky's theorem** and GMM-WGHT assumption, we have that:

$$n^{1/2}(\hat{\theta}_{GMM} - \theta) = - [D^T \mathbf{W}_n D]^{-1} \times D^T \mathbf{W}_n n^{1/2} \bar{g}_n(\theta) + o_p(1) \xrightarrow{D} - [D^T \mathbf{W} D]^{-1} D^T \mathbf{W} \cdot \mathcal{N}(0, S)$$

where we have that:

$$[D^T W D]^{-1} D^T W \cdot \mathcal{N}(\mathbf{0}, \mathbf{S}) = \mathcal{N}(\mathbf{0}, (D^T W D)^{-1} D^T W S W D (D^T W D)^{-1}).$$

Notice that in these past steps, we assumed that

$$D^T W D$$

was invertible. Well, we know that $\mathbf{W}$ is a positive-definite matrix by assumption. Therefore, we have that:

$$\text{Rank}(D^T W D) = \text{Rank}(D)$$

as $\mathbf{W}$ is a positive definite matrix. Therefore, in order for $D^T W D$ to be invertible, we require that the matrix D is full-rank. However, this was our assumption in GMM-JAC! Therefore,

$$\text{Rank}(D) = \text{Rank}[D_\theta \mathbb{E}[g(w_i; \theta)]] = K$$

where $D_\theta \mathbb{E}[g(w_i; \theta)]$ is a $L \times K$ matrix.

To conclude, we have shown that:

$$n^{1/2}(\hat{\theta}_{GMM} - \theta) \xrightarrow{D} \mathcal{N}(\mathbf{0}, \mathbf{V}_W)$$

where the variance matrix is given by:

$$\mathbf{V}_W = (D^T W D)^{-1} D^T W S W D (D^T W D)^{-1}.$$

□

In practice, the Jacobian D and the variance-of-moments matrix $\mathbf{S}$ are unknown because they depend on the true population parameter θ . We therefore replace them with estimators evaluated at $\hat{\theta}_{GMM}$:

$$\hat{D} \equiv D_\theta \bar{g}_n(\hat{\theta}_{GMM}) = \frac{1}{n} \sum_{i=1}^n D_\theta g(w_i; \hat{\theta}_{GMM}). \quad (293)$$

$$\hat{S}_n \equiv \frac{1}{n} \sum_{i=1}^n g(w_i; \hat{\theta}_{GMM}) g(w_i; \hat{\theta}_{GMM})^T - \bar{g}_n(\hat{\theta}_{GMM}) \bar{g}_n(\hat{\theta}_{GMM})^T \quad (294)$$

11.4 Identification

We now want to look more into the rank of the Jacobian assumption in [GMM-Jacobian]. We can in fact see that the rank of the Jacobian of the population moments are in fact connected to identification.

Under the GMM-ID assumption, we have that $\mathbb{E}[g(w_i; \tilde{\theta})] = \mathbf{0}$ if and only if $\tilde{\theta} = \theta$. We say that θ is **globally identified** as there is no other parameter in the model that can fulfill this condition. The distribution of the data uniquely pins down the parameter of interest.

Under the GMM-JAC assumption, we assumed that the rank of the Jacobian of the population moment condition was full column rank:

$$D \equiv D_\theta \mathbb{E}[g(w_i; \theta)]$$

was rank K . This is a sufficient local rank condition and is necessary for the usual regular, asymptotically normal GMM theory, although it is not necessary for identification in every nonlinear model. We can view this condition as a local

identification property. To see why, we say that θ is **locally identified** if there exists an $\epsilon > 0$ such that

$$\mathbb{E}[g(w_i; \tilde{\theta})] \neq 0$$

for all $0 < \|\tilde{\theta} - \theta\| < \epsilon$. This says that there is no other $\tilde{\theta}$ that is a solution to the moment conditions besides the true population parameter θ within a neighborhood. To see how this comes about, let us take a Taylor expansion of the moment condition around a small neighborhood of the true population parameter:

$$\mathbb{E}[g(w_i; \tilde{\theta})] \approx \mathbb{E}[g(w_i; \theta)] + D(\tilde{\theta} - \theta)$$

The first term is zero by definition $\mathbb{E}[g(w_i; \theta)] = 0$. For the second term, we use the linear algebra result that if D has **full column rank** then the **only** way for $D(\tilde{\theta} - \theta) = 0$ is if the vector $\tilde{\theta} - \theta = \mathbf{0}$ which only holds if $\tilde{\theta} = \theta$. Therefore, we have that the moment conditions:

$$\mathbb{E}[g(w_i; \tilde{\theta})] \approx 0$$

in a small neighborhood around the true population parameter θ if D is full rank and $\tilde{\theta} = \theta$.

The issue with local identification is that it does not preclude other parameters that lie outside of the neighborhood that satisfies the moment conditions.

Proposition 363 (Global and Local Identification for GMM)

In the GMM framework, a parameter is **globally identified** under the [GMM-ID] assumption. A parameter is **locally identified** if [GMM-JAC] assumption is satisfied.

In practice, we can analyse whether local identification holds by considering the rank of $D_{\theta} \bar{g}_n(\hat{\theta}_{GMM})$, as this is a necessary condition for the asymptotic normality of the GMM estimator. If this matrix is not full rank, we have big issues.

11.4.1 Identification in Linear Model with Instrumental Variables

We apply our ideas of local identification to the instrumental variables framework. Recall that the [GMM-JAC] assumption needed for local identification was that the rank of the Jacobian of the population moment conditions:

$$\text{Rank}(D) = \text{Rank}(D_{\theta} \mathbb{E}[g(w_i; \theta)]) = K$$

is full column rank. In the IV model, recall that $g(w_i; \beta) = z_i(y_i - x_i^T \beta)$. Therefore, the population moment condition is:

$$\mathbb{E}[g(w_i; \beta)] = \mathbb{E}[z_i(y_i - x_i^T \beta)]$$

Then, the Jacobian of this is:

$$D = D_{\beta} \mathbb{E}[z_i(y_i - x_i^T \beta)] = -\mathbb{E}[z_i x_i^T]$$

We require that $\mathbb{E}[z_i x_i^T]$ has rank K for local identification. However, recall that this also coincided with the global identification required in IV models as we saw in Proposition 339! Therefore, in the IV framework, global and local identification are equivalent to each other. This is due to the fact that the moment conditions are linear.

Theorem 364 (Global and Local Identification in IV Framework)

In the instrumental variables framework, global and local identification are equivalent.

11.5 Asymptotic Efficiency

We have shown that asymptotic normality holds for GMM estimators. However, in order to do inference, we require more efficient regarding the variance of the estimator.

First, we have already shown that under [GMM-WGHT], the weight matrix

$$\mathbf{W}_n \xrightarrow{P} \mathbf{W}.$$

From this, the asymptotic variance of the GMM estimator $\hat{\theta}_{GMM}(\mathbf{W}_n)$ was shown to be:

$$\mathbf{V}_W = (\mathbf{D}^T \mathbf{W} \mathbf{D})^{-1} \mathbf{D}^T \mathbf{W} \mathbf{S} \mathbf{W} \mathbf{D} (\mathbf{D}^T \mathbf{W} \mathbf{D})^{-1}$$

as we saw in Theorem 358 whereby $\mathbf{D}$ is the Jacobian of the moments at the true population parameters

$$\mathbf{D} \equiv \mathbf{D}_\theta \mathbb{E}[g(w_i; \theta)]$$

and $\mathbf{S}$ is the variance of the moments

$$\mathbf{S} = \mathbb{E}[g(w_i; \theta)g(w_i; \theta)^T].$$

The question we wish to ask is what is the weight matrix $\mathbf{W}$ that minimises the variance $\mathbf{V}_W$.

Theorem 365 (Optimal Weight Matrix for GMM)

Setting the weight matrix as the inverse of the population covariance matrix of the moments,

$$\mathbf{W} = \mathbf{S}^{-1},$$

minimises the variance of the GMM estimator in the sense that

$$\mathbf{V}_W - \mathbf{V}_{S^{-1}}$$

is positive semidefinite for any other weight matrix $\mathbf{W}$.

The variance of the GMM estimator under this weight matrix $\mathbf{W} = \mathbf{S}^{-1}$ is:

$$\mathbf{V}_{S^{-1}} = (\mathbf{D}^T \mathbf{S}^{-1} \mathbf{D})^{-1} \tag{295}$$

To interpret what $\mathbf{V}_W - \mathbf{V}_{S^{-1}}$ is doing, the optimal weighting scheme downweights noisy linear combinations of moments and also accounts for covariance across the moments. This ensures that the GMM estimator $\hat{\theta}_{GMM}$ is informed most strongly by moments that are estimated more precisely.

However, one issue is that we don't know what $\mathbf{S} = \mathbb{E}[g(w_i; \theta)g(w_i; \theta)^T]$ is in practice since we do not know the true population parameter θ . There are two ways around this.

The first way around this is to compute the **two-step GMM** estimator. In the first step, estimate θ using any consistent initial positive-definite weight matrix, often the identity matrix. Then estimate $\mathbf{S}$ from the moments evaluated at this first-step estimate. Finally, re-estimate θ using $\hat{\mathbf{S}}^{-1}$ as the weight matrix.

The second way is to use **continuously-updated GMM**. Here, for each conjectured value $\tilde{\theta}$, we calculate the variance of sample moments at that conjecture $\tilde{\theta}$, which we denote by $S_n(\tilde{\theta})$. The continuously updated estimator is then the quadratic form that is allowed to depend on this conjecture and weight matrix now depends on $\tilde{\theta}$.

11.5.1 Efficiency in Linear IV Model

We now analyse the efficiency of the GMM estimator in the instrumental variable framework.

Theorem 366 (Optimal Weight Matrix for Linear IV Model Under Homoskedasticity)

Under homoskedasticity, the optimal weight matrix for the linear IV model is:

$$\mathbf{W}_n^* = \left(\frac{1}{n} \sum_{i=1}^n \mathbf{z}_i \mathbf{z}_i^T \right)^{-1} = n \left(\mathbf{Z}^T \mathbf{Z} \right)^{-1} \propto \left(\mathbf{Z}^T \mathbf{Z} \right)^{-1} \quad (296)$$

Proof. First, we write out our sample moment conditions:

$$\bar{g}_n(\tilde{\boldsymbol{\beta}}) = \frac{1}{n} \sum_{i=1}^n g(w_i; \tilde{\boldsymbol{\beta}}) = \frac{1}{n} \sum_{i=1}^n \mathbf{z}_i (y_i - \mathbf{x}_i^T \tilde{\boldsymbol{\beta}}) \stackrel{(1)}{=} \frac{1}{n} \sum_{i=1}^n \mathbf{z}_i u_i$$

where $\stackrel{(1)}{=}$ holds if $\tilde{\boldsymbol{\beta}} = \boldsymbol{\beta}$ which is the true population parameter.

Therefore, the variance of the sample moments at the true population parameters is then given by:

$$\mathbf{S} = \mathbb{E}[g(w_i; \boldsymbol{\beta}) g(w_i; \boldsymbol{\beta})^T] = \mathbb{E}[u_i^2 \mathbf{z}_i \mathbf{z}_i^T] \stackrel{(1)}{=} \sigma_u^2 \mathbb{E}[\mathbf{z}_i \mathbf{z}_i^T]$$

where $\stackrel{(1)}{=}$ requires the assumption of **homoskedasticity**.

We require an estimator of the variance $\mathbf{S}$, whereby we can use:

$$\hat{\mathbf{S}}_n \equiv \hat{\sigma}_n^2 \left(\frac{1}{n} \sum_{i=1}^n \mathbf{z}_i \mathbf{z}_i^T \right)$$

whereby

$$\hat{\sigma}_n^2 = \frac{1}{n} \sum_{i=1}^n \hat{u}_i^2$$

is the residual-based estimator from 2SLS, which gives us a consistent estimator of σ^2 .

Then, recall that the optimal weight matrix is the inverse of the variance but we now plug in our estimated variance:

$$\mathbf{W}_n^* = \hat{\mathbf{S}}_n^{-1} = \hat{\sigma}_n^{-2} \left(\frac{1}{n} \sum_{i=1}^n \mathbf{z}_i \mathbf{z}_i^T \right)^{-1}$$

Now, notice that scaling the weight matrix $\mathbf{W}_n^*$ has no effect on the optimisation routine of GMM:

$$\min_{\tilde{\boldsymbol{\beta}}} \bar{g}_n(\tilde{\boldsymbol{\beta}})^T \mathbf{W}_n^* \bar{g}_n(\tilde{\boldsymbol{\beta}})$$

since it is only the relative importance of weights on the moments that matters, not their absolute levels. Therefore, we can drop the scalar factor $\hat{\sigma}_n^{-2}$ from the optimal weight matrix:

$$\mathbf{W}_n^* = \left(\frac{1}{n} \sum_{i=1}^n \mathbf{z}_i \mathbf{z}_i^T \right)^{-1} = n \left(\mathbf{Z}^T \mathbf{Z} \right)^{-1} \propto \left(\mathbf{Z}^T \mathbf{Z} \right)^{-1}$$

□

Now, recall that we have shown that in Proposition 350, such a weight matrix gives us a two-stage least squares estimator! Hence, under homoskedasticity, the 2SLS estimator is the asymptotically efficient GMM estimator.

Under **heteroskedasticity**, we have that the variance matrix

$$\mathbf{S} = \mathbb{E}[u_i^2 \mathbf{z}_i \mathbf{z}_i^T]$$

can be estimated by:

$$\hat{\mathbf{S}}_n = \frac{1}{n} \sum_{i=1}^n \hat{u}_i^2 \mathbf{z}_i \mathbf{z}_i^T$$

The weight matrix is then set to be:

$$\mathbf{W}_n = (\hat{\mathbf{S}}_n)^{-1}$$

We can then compute the GMM estimator under this weight matrix.

11.6 Test For Overidentifying Restrictions

We are now interested in looking at testing overidentifying restrictions in the GMM framework. This is analogous to instrument exogeneity testing as we shall soon see.

Definition 367 (Hypothesis Test for Overidentifying Restrictions)

The hypothesis test for overidentifying restrictions is:

$$H_0 : \mathbb{E}[g(w_i; \boldsymbol{\theta})] = 0$$

against

$$H_1 : \mathbb{E}[g(w_i; \boldsymbol{\theta})] \neq 0$$

Notice that this is similar to testing instrument exogeneity:

$$H_0 : \mathbb{E}[\mathbf{z}_i(y_i - \mathbf{x}_i^T \boldsymbol{\beta})] = 0.$$

We base our test for overidentifying restrictions on the GMM criterion function we defined in Definition 345:

$$Q_n(\tilde{\boldsymbol{\theta}}) = \bar{\mathbf{g}}_n(\tilde{\boldsymbol{\theta}})^T \mathbf{W}_n \bar{\mathbf{g}}_n(\tilde{\boldsymbol{\theta}}) \quad (297)$$

whereby $\mathbf{W}_n$ is a positive definite $L \times L$ weight matrix and $\bar{\mathbf{g}}_n(\tilde{\boldsymbol{\theta}})$ are the sample moment conditions. Recall that Q_n is a scalar!

Theorem 368 (Limiting Distribution for Test for Overidentifying Restriction in GMM)

Let Q_n be the criterion function. Then, if the optimal weight matrix $\mathbf{W}_n \xrightarrow{P} \mathbf{S}^{-1}$, the limiting distribution of the rescaled GMM criterion function is:

$$nQ_n(\hat{\boldsymbol{\theta}}_{GMM}) \xrightarrow{D} \chi^2(L - K) \quad (298)$$

where $L - K$ is the number of overidentified restrictions.

Proof. First, look at the sample moments $\bar{\mathbf{g}}_n(\tilde{\boldsymbol{\theta}})$. We shall examine them under the true population parameters $\bar{\mathbf{g}}_n(\boldsymbol{\theta})$.

Theorem 369 (Asymptotic Normality of Sample Moments at True Parameter)

The sample moments under the true population parameters θ is asymptotically normal:

$$n^{1/2}\bar{g}_n(\theta) \xrightarrow{D} \mathcal{N}(\mathbf{0}, \mathbf{S}) \quad (299)$$

where $\mathbf{S}$ is the variance of the moments $\mathbf{S} = \mathbb{E}[g(w_i; \theta)g(w_i; \theta)^T]$.

Proof. The sample moment averages at the true parameter is given by:

$$\bar{g}_n(\theta) = \frac{1}{n} \sum_{i=1}^n g(w_i; \theta)$$

Clearly, by definition, the moments have population mean 0 at the true parameter

$$\mathbb{E}[g(w_i; \theta)] = \mathbf{0}$$

by the GMM-ID assumption. Therefore, we have a sample average with mean 0. By assumption, w_i are i.i.d. We can divide by $\sqrt{n}$ and invoke the **central limit theorem**

$$n^{1/2}\bar{g}_n(\theta) = \frac{1}{n^{1/2}} \sum_{i=1}^n g(w_i; \theta) \xrightarrow{D} \mathcal{N}(\mathbf{0}, \mathbf{S})$$

whereby $\mathbf{S} \equiv \mathbb{E}[g(w_i; \theta)g(w_i; \theta)^T]$ is the variance-covariance matrix of the sample moments as the expectation is 0. □

We can rewrite the result we have just shown as:

$$n^{1/2}\bar{g}_n(\theta) \xrightarrow{D} \mathcal{N}(\mathbf{0}, \mathbf{S}) = \mathbf{S}^{1/2}\boldsymbol{\eta}$$

whereby $\boldsymbol{\eta} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_L)$ is a multivariate normal distribution and $\mathbf{S}^{1/2}$ is the square root of the variance matrix.

We now assume that we know the true value θ . Now, we multiply the GMM criterion function Eq. (297) by n and apply **Slutsky's theorem**, the GMM criterion at the true population parameter θ is:

$$nQ_n(\theta) = (\sqrt{n}\bar{g}_n(\theta))^T \mathbf{W}_n(\sqrt{n}\bar{g}_n(\theta)) \xrightarrow{D} \boldsymbol{\eta}^T \mathbf{S}^{1/2} \mathbf{W} \mathbf{S}^{1/2} \boldsymbol{\eta}$$

under the null hypothesis $H_0 : \mathbb{E}[g(w_i; \theta)] = \mathbf{0}$.

Now, suppose that the optimal weight matrix is chosen: $\mathbf{W} = \mathbf{S}^{-1}$. Then, our limiting expression for the rescaled criterion function is:

$$nQ_n(\theta) \xrightarrow{D} \boldsymbol{\eta}^T \mathbf{S}^{1/2} \mathbf{S}^{-1} \mathbf{S}^{1/2} \boldsymbol{\eta} = \boldsymbol{\eta}^T \boldsymbol{\eta}$$

Now, recall that $\boldsymbol{\eta} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_L)$ is a multivariate standard normal vector. Its squared Euclidean norm therefore has a χ^2 distribution:

$$\boldsymbol{\eta}^T \boldsymbol{\eta} \sim \chi^2(L)$$

However, we don't know the true value θ , so we base our test on the GMM estimator $\hat{\theta}_{GMM}$, which is a consistent estimator for the true parameter θ under the null H_0 . First, recall that the first order condition for the GMM criterion:

$$\mathbf{0} = \nabla_{\hat{\theta}} Q_n(\hat{\theta}) = \left[D_{\hat{\theta}} \bar{g}_n(\hat{\theta}_{GMM}) \right]^T \mathbf{W}_n \bar{g}_n(\hat{\theta}_{GMM})$$

This expression is a $K \times 1$ vector because $D_{\hat{\theta}} \bar{g}_n(\hat{\theta}_{GMM})$ is an $L \times K$ matrix, $\mathbf{W}_n$ is an $L \times L$ matrix, and $\bar{g}_n(\hat{\theta}_{GMM})$

is an $L \times 1$ vector. By what we have shown before, $K \leq L$. Regardless of whether H_0 is true, the first-order conditions set K linear combinations of the sample moments to zero by construction.

If we have as many moments as parameters ($K = L$), then all moment conditions can be set to 0 and we have no overidentifying restrictions to test. When $L > K$, estimation uses K independent directions in the moments, leaving $L - K$ independent overidentifying restrictions. Therefore, the test has $L - K$ degrees of freedom. Hence,

$$nQ_n(\hat{\theta}_{GMM}) \xrightarrow{D} \chi^2(L - K).$$

□

11.7 Hypothesis Testing

We are now interested in testing hypotheses involving linear or nonlinear restrictions.

11.7.1 Linear Testing

We have shown that the GMM estimator is asymptotically normal:

$$n^{1/2}[\hat{\theta}_{GMM}(W_n) - \theta] \xrightarrow{D} \mathcal{N}(0, V_W)$$

where the variance matrix is given by:

$$V_W = (D^T W D)^{-1} D^T W S W D (D^T W D)^{-1}.$$

Definition 370 (Linear Restrictions)

The null hypothesis of testing linear restrictions is given by:

$$H_0 : R\theta = \rho \quad H_1 : R\theta \neq \rho$$

where R is a $r \times K$ matrix, ρ is an $r \times 1$ vector, and r is the number of restrictions, with $r \leq K$.

Under asymptotic normality, we can construct the Wald test.

Definition 371 (Wald Statistic)

The **Wald Statistic** is given by:

$$W_n = n(\hat{R}\hat{\theta}_{GMM} - \rho)^T (\hat{R}\hat{V}_n\hat{R}^T)^{-1} (\hat{R}\hat{\theta}_{GMM} - \rho). \quad (300)$$

With the Wald statistic, we can derive the asymptotic distribution needed to conduct hypothesis testing.

Theorem 372 (Wald Test)

Let W_n be the Wald Statistic. The limiting distribution of the Wald Statistic is a chi-squared distribution:

$$W_n \xrightarrow{D} \chi^2(r) \quad (301)$$

where r is the number of restrictions made in the null hypothesis.

Proof. Recall the Wald statistic

$$W_n = n(\mathbf{R}\hat{\boldsymbol{\theta}}_{GMM} - \boldsymbol{\rho})^T (\mathbf{R}\hat{\mathbf{V}}_n \mathbf{R}^T)^{-1} (\mathbf{R}\hat{\boldsymbol{\theta}}_{GMM} - \boldsymbol{\rho})$$

First, by consistency, we have that the middle expression is:

$$\mathbf{R}\hat{\mathbf{V}}_n \mathbf{R}^T \xrightarrow{P} \mathbf{R}\mathbf{V}\mathbf{R}^T.$$

We now look at the *bread* of the Wald statistic and see we can rewrite it:

$$\begin{aligned} n^{1/2}(\mathbf{R}\hat{\boldsymbol{\theta}}_{GMM} - \boldsymbol{\rho}) &=_{(1)} n^{1/2}(\mathbf{R}\hat{\boldsymbol{\theta}}_{GMM} - \mathbf{R}\boldsymbol{\theta}) \\ &= n^{1/2}(\mathbf{R}(\hat{\boldsymbol{\theta}}_{GMM} - \boldsymbol{\theta})) \xrightarrow{D} \mathbf{R}\mathcal{N}(\mathbf{0}, \mathbf{V}) \\ &= \mathcal{N}(\mathbf{0}, \mathbf{R}\mathbf{V}\mathbf{R}^T) \\ &= (\mathbf{R}\mathbf{V}\mathbf{R}^T)^{1/2} \boldsymbol{\eta} \quad \text{where } \boldsymbol{\eta} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_r) \end{aligned}$$

whereby $=_{(1)}$ holds under the null hypothesis of the linear restriction and we invoked the asymptotic normality of the GMM estimator.

So going back to our Wald Statistic, we combine our results with Slutsky's theorem to get:

$$\begin{aligned} W_n &= n(\mathbf{R}\hat{\boldsymbol{\theta}}_{GMM} - \boldsymbol{\rho})^T (\mathbf{R}\hat{\mathbf{V}}_n \mathbf{R}^T)^{-1} (\mathbf{R}\hat{\boldsymbol{\theta}}_{GMM} - \boldsymbol{\rho}) \\ &= (\sqrt{n}(\mathbf{R}\hat{\boldsymbol{\theta}}_{GMM} - \boldsymbol{\rho}))^T (\mathbf{R}\hat{\mathbf{V}}_n \mathbf{R}^T)^{-1} (\sqrt{n}(\mathbf{R}\hat{\boldsymbol{\theta}}_{GMM} - \boldsymbol{\rho})) \\ &\xrightarrow{D} [(\mathbf{R}\mathbf{V}\mathbf{R}^T)^{1/2} \boldsymbol{\eta}]^T (\mathbf{R}\mathbf{V}\mathbf{R}^T)^{-1} [(\mathbf{R}\mathbf{V}\mathbf{R}^T)^{1/2} \boldsymbol{\eta}] \\ &= \boldsymbol{\eta}^T (\mathbf{R}\mathbf{V}\mathbf{R}^T)^{1/2} (\mathbf{R}\mathbf{V}\mathbf{R}^T)^{-1} (\mathbf{R}\mathbf{V}\mathbf{R}^T)^{1/2} \boldsymbol{\eta} \\ &=_{(1)} \boldsymbol{\eta}^T \boldsymbol{\eta} \sim \chi^2(r) \end{aligned}$$

where $=_{(1)}$ uses the fact that the squared Euclidean norm of an r -dimensional standard normal vector has a $\chi^2(r)$ distribution. $\square$

Therefore, for linear restrictions, we construct a Wald statistic and compare its value to a $\chi^2(r)$ distribution and reject the null hypothesis if the Wald statistic is greater than the specified critical value for the $\chi^2(r)$ distribution.

11.7.2 Nonlinear Testing

We now look at testing nonlinear restrictions. First, we change notation and let d_r be the number of restrictions we are testing and r be the restrictions themselves. Now, we assume that the restrictions $r : \Theta \rightarrow \mathbb{R}^{d_r}$ are represented by a differentiable function with the Jacobian of the restrictions defined as:

$$\mathbf{R} \equiv D_{\boldsymbol{\theta}} r(\boldsymbol{\theta})$$

whereby $D_{\boldsymbol{\theta}} r(\boldsymbol{\theta})$ is a $d_r \times K$ matrix having rank $d_r \leq K$.

Definition 373 (Hypothesis Test)

The hypothesis test for nonlinear restrictions are the hypotheses:

$$H_0 : r(\boldsymbol{\theta}) = \boldsymbol{\rho} \quad \text{vs} \quad H_1 : r(\boldsymbol{\theta}) \neq \boldsymbol{\rho}$$

We can now define the Wald statistic which shall be a quadratic form of the restrictions evaluated at the consistent estimators of the true population parameter.

Definition 374 (Wald Statistic)

The Wald Statistic under nonlinear restrictions is given by:

$$W_n = n(r(\hat{\theta}_{GMM}) - \rho)^T (R\hat{V}_n R^T)^{-1} (r(\hat{\theta}_{GMM}) - \rho) \quad (302)$$

We now state the asymptotic distribution for the Wald statistic under nonlinear restrictions.

Theorem 375 (Limiting Distribution of Wald Statistic for Nonlinear Restrictions)

The Wald statistic under nonlinear restrictions has a limiting chi-squared distribution:

$$W_n \xrightarrow{D} \chi^2(d_r). \quad (303)$$

Proof. First, we rewrite the Wald Statistic:

$$\begin{aligned} W_n &= n(r(\hat{\theta}_{GMM}) - \rho)^T (R\hat{V}_n R^T)^{-1} (r(\hat{\theta}_{GMM}) - \rho) \\ &= \left(n^{1/2}(r(\hat{\theta}_{GMM}) - \rho) \right)^T (R\hat{V}_n R^T)^{-1} \left(\sqrt{n}(r(\hat{\theta}_{GMM}) - \rho) \right) \end{aligned}$$

The *meat* of our sandwich uses the fact that $\hat{V}$ is consistent to show:

$$R\hat{V}_n R^T \xrightarrow{P} RVR^T.$$

We now look at the bread of our sandwich.

$$n^{1/2}(r(\hat{\theta}_{GMM}) - \rho) = n^{1/2}(r(\hat{\theta}_{GMM}) - r(\theta))$$

whereby we imposed the null hypothesis $r(\theta) = \rho$. Now, we know that the GMM estimator is asymptotically normal:

$$n^{1/2}(\hat{\theta}_{GMM} - \theta) \xrightarrow{D} \mathcal{N}(\mathbf{0}, \mathbf{V})$$

and we are interested in the distribution of a function $r(\cdot)$ of $\hat{\theta}_{GMM}$. We can apply a matrix version of the **delta method** to derive the limiting distribution of $r(\hat{\theta}_{GMM})$ as:

$$n^{1/2}(r(\hat{\theta}_{GMM}) - r(\theta)) \xrightarrow{D} [D_\theta r(\theta)] \mathcal{N}(\mathbf{0}, \mathbf{V})$$

However, by definition, we said that:

$$\mathbf{R} \equiv D_\theta r(\theta)$$

so we have that:

$$[D_\theta r(\theta)] \mathcal{N}(\mathbf{0}, \mathbf{V}) = \mathbf{R} \mathcal{N}(\mathbf{0}, \mathbf{V}) = \mathcal{N}(\mathbf{0}, \mathbf{RVR}^T) = (\mathbf{RVR}^T)^{1/2} \boldsymbol{\eta}$$

whereby $\boldsymbol{\eta} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_{d_r})$. We can combine our two results and apply Slutsky's theorem to conclude that:

$$\begin{aligned}
W_n &= \left(n^{1/2}(r(\hat{\theta}_{GMM}) - \rho) \right)^T (R\hat{V}_n R^T)^{-1} \left(\sqrt{n}(r(\hat{\theta}_{GMM}) - \rho) \right) \\
&\xrightarrow{D} \left((RV R^T)^{1/2} \eta \right)^T \left(RV R^T \right)^{-1} \left((RV R^T)^{1/2} \eta \right) \\
&= \eta^T \eta \sim \chi^2(d_r)
\end{aligned}$$

whereby we used the fact that the squared Euclidean norm of a d_r -dimensional standard normal vector has a chi-squared distribution. $\square$

In practice, we don't know what R is as it can vary with the true population value θ . Therefore, we can construct a consistent estimator of R with:

$$\hat{R}_n = D_{\theta} r(\hat{\theta}_{GMM}) \xrightarrow{P} D_{\theta} r(\theta) = R$$

using the GMM estimator $\hat{\theta}_{GMM}$. From this, we can plug this into our Wald Statistic:

$$W_n = n(r(\hat{\theta}_{GMM}) - \rho)^T (\hat{R}_n \hat{V}_n \hat{R}_n^T)^{-1} (r(\hat{\theta}_{GMM}) - \rho)$$

and this can be shown to still converge to an asymptotic χ^2 distribution.

11.8 Criterion-Based Test

There are two issues with the Wald statistic. The first is that it is not invariant to equivalent representations of the restrictions, so restrictions that are essentially equivalent can give drastically different test results. Second, the Jacobian of the restrictions may fail to have full rank. The criterion-based test addresses the invariance issue, although a rank failure can still invalidate the usual chi-squared approximation.

We can construct something called the **QLR test** which is similar to the likelihood ratio test.

First, define the null hypotheses:

$$H_0 : r(\theta) = \rho \quad H_1 : r(\theta) \neq \rho$$

Recall that the **GMM criterion** is given by:

$$Q_n(\tilde{\theta}) \equiv \bar{g}_n(\tilde{\theta})^T \hat{S}_n^{-1} \bar{g}_n(\tilde{\theta}).$$

The QLR test works by evaluating the fit of the model under the null hypothesis and not under the null hypothesis. That is, we compare the criterion function values between the restricted and unrestricted model.

Definition 376 (QLR Test)

Define a subset of the parameter space:

$$\Theta_p \equiv \{\theta \in \Theta : r(\theta) = \rho\}$$

as the parameters that satisfy the null hypothesis. Then, the QLR statistic is given by:

$$QLR_n \equiv n \left[\min_{\tilde{\theta} \in \Theta_p} Q_n(\tilde{\theta}) - \min_{\tilde{\theta} \in \Theta} Q_n(\tilde{\theta}) \right] \geq 0$$

where the first term is the criterion function under the restricted model and the second term is under the unrestricted model.

We can also derive the limiting distribution of this test.

Proposition 377 (Limiting Distribution of QLR Test)

The limiting distribution of the QLR test is:

$$QLR_n \xrightarrow{D} \chi^2(d_r) \quad (304)$$

where d_r is the number of restrictions.

Whilst the Wald statistic and QLR test statistic are asymptotically equivalent under standard regularity conditions, the QLR test is often preferable because it is invariant to equivalent reparameterisations of the restrictions.

12 Time Series

12.1 Introduction

The big issue in time series data is that observations and error terms may be serially dependent, unlike in the simplest cross-sectional setting. Lagged variables can help us model temporal dependence in the conditional mean, whilst the remaining dependence in the errors must also be handled by the specification or by appropriate inference.

12.2 Cross Sectional Data Review

We now review concepts we have previously seen in cross-sectional data.

12.2.1 Ordinary Least Squares

We first revisit regressions that we saw in the cross sectional case. In particular, it is very important to understand how the Frisch-Waugh-Lovell theorem works as that is the key behind how other concepts in time series operate.

We assume a statistical model in matrix notation given by:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}.$$

The **OLS estimator** is given by:

$$\hat{\boldsymbol{\beta}} = \left(\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^T \right)^{-1} \left(\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i y_i \right) = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$$

whereby we require the assumption that $\mathbf{X}^T \mathbf{X}$ is invertible.

The **least squares fitted values** is given by:

$$\hat{\mathbf{y}} = \mathbf{X}\hat{\boldsymbol{\beta}} = \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} = \mathbf{P}\mathbf{y}$$

whereby we defined the **projection matrix**:

$$\mathbf{P} \equiv \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T.$$

The **least squares residuals** is given by:

$$\hat{\mathbf{u}} = \mathbf{y} - \hat{\mathbf{y}} = (\mathbf{I} - \mathbf{P})\mathbf{y} = \mathbf{M}\mathbf{y}$$

where we defined the **annihilator matrix**:

$$\mathbf{M} = \mathbf{I} - \mathbf{P}.$$

We have that $\mathbf{P}, \mathbf{M}$ are idempotent and symmetric matrices, and are orthogonal to each other. The **residual variance** can be estimated by:

$$s^2 = \frac{1}{n-k} \hat{\mathbf{u}}^T \hat{\mathbf{u}}.$$

12.2.2 Maximum Likelihood Estimation

We now revisit maximum likelihood estimation which we will also need. We state the maximised likelihood function.

Proposition 378 (Likelihood Function under Normality)

Suppose that y_i, x_i are i.i.d and $u_i \sim \mathcal{N}(0, \sigma^2)$. Then, the likelihood function is given by:

$$\hat{\ell} = -\frac{n}{2} \log(2\pi e \hat{\sigma}^2) \quad (305)$$

whereby:

$$\hat{\sigma}^2 = \frac{1}{n} \hat{u}^T \hat{u}.$$

Proof. The log-likelihood function is given by:

$$\ell_{y|x}(\boldsymbol{\beta}, \sigma^2) = -\frac{n}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} \sum_{i=1}^n (y_i - \mathbf{x}_i^T \boldsymbol{\beta})^2 \quad (306)$$

The maximum likelihood estimation for $\boldsymbol{\beta}$ clearly involves minimising the second term as that is the term that depends on $\boldsymbol{\beta}$. Doing so gives us:

$$\hat{\boldsymbol{\beta}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}.$$

The estimate for the variance is:

$$\hat{\sigma}^2 = \frac{1}{n} \hat{\mathbf{u}}^T \hat{\mathbf{u}}.$$

We plug these back into Eq. (306):

$$\begin{aligned} \hat{\ell} &= -\frac{n}{2} \log(2\pi \hat{\sigma}^2) - \frac{1}{2\hat{\sigma}^2} \sum_{i=1}^n (\hat{u}_i)^2 \\ &= -\frac{n}{2} \log(2\pi \hat{\sigma}^2) - \frac{1}{2\hat{\sigma}^2} n \hat{\sigma}^2 \\ &= -\frac{n}{2} \log(2\pi \hat{\sigma}^2) - \frac{n}{2} \\ &= -\frac{n}{2} \log(2\pi \hat{\sigma}^2) - \frac{n}{2} \log(e) \\ &= -\frac{n}{2} \left[\log(2\pi \hat{\sigma}^2) + \log(e) \right] \\ &= -\frac{n}{2} \log(2\pi e \hat{\sigma}^2). \end{aligned}$$

□

This is useful because we now have an analytical expression for the likelihood ratio test. Suppose we had the model:

$$\mathbf{y} = \mathbf{X}_1 \boldsymbol{\beta}_1 + \mathbf{X}_2 \boldsymbol{\beta}_2 + \mathbf{u}$$

and we wish to test $\beta_2 = 0$. We can conduct a log-likelihood ratio test instead of a t-test.

Proposition 379 (Log-Likelihood Ratio Test)

The log-likelihood ratio test is given by:

$$LR_{\beta_2=0} = n \log \left(\frac{\hat{\sigma}_{Restricted}^2}{\hat{\sigma}_{Unrestricted}^2} \right) \geq 0.$$

Proof. The log-likelihood ratio test is given by:

$$LR_{\beta_2=0} = -2 \left\{ \max_{\beta_1, \sigma^2} \ell(\beta_1, \mathbf{0}, \sigma^2) - \max_{\beta_1, \beta_2, \sigma^2} \ell(\beta_1, \beta_2, \sigma^2) \right\} \geq 0$$

We then can just plug in our expression for the maximised log-likelihood $\hat{\ell} = -\frac{n}{2} \log(2\pi e \hat{\sigma}^2)$ to get the desired result. $\square$

12.2.3 Partial Least Squares

We are now interested in **partial least squares estimators**. The first major result we state is the Frisch-Waugh-Lovell theorem.

Theorem 380 (Frisch-Waugh-Lovell Theorem)

Suppose we had the model:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}$$

Then suppose we partitioned the design matrix:

$$\mathbf{X} = \begin{bmatrix} \mathbf{X}_1 & \mathbf{X}_2 \end{bmatrix}$$

whereby $\mathbf{X}_1$ is a $n \times k_1$ matrix and $\mathbf{X}_2$ is a $n \times k_2$ matrix. Likewise, partition the coefficients:

$$\boldsymbol{\beta} = \begin{bmatrix} \boldsymbol{\beta}_1 \\ \boldsymbol{\beta}_2 \end{bmatrix}$$

whereby $\boldsymbol{\beta}_1$ is a $k_1 \times 1$ vector and $\boldsymbol{\beta}_2$ is a $k_2 \times 1$ vector.

The **Frisch-Waugh-Lovell theorem** says that:

$$\hat{\boldsymbol{\beta}}_2 = (\hat{\mathbf{u}}_{2,1}^T \hat{\mathbf{u}}_{2,1})^{-1} \hat{\mathbf{u}}_{2,1}^T \hat{\mathbf{u}}_{y,1}.$$

where $\hat{\mathbf{u}}_{y,1}$ are the residuals from regressing $\mathbf{y}$ on $\mathbf{X}_1$ and $\hat{\mathbf{u}}_{2,1}$ are the residuals from regression $\mathbf{X}_2$ on $\mathbf{X}_1$.

Proof. First, let us partition our design matrix:

$$\mathbf{X} = \begin{bmatrix} \mathbf{X}_1 & \mathbf{X}_2 \end{bmatrix}$$

whereby $\mathbf{X}_1$ is a $n \times k_1$ matrix and $\mathbf{X}_2$ is a $n \times k_2$ matrix. Likewise, we partition the coefficients:

$$\boldsymbol{\beta} = \begin{bmatrix} \boldsymbol{\beta}_1 \\ \boldsymbol{\beta}_2 \end{bmatrix}$$

whereby $\boldsymbol{\beta}_1$ is a $k_1 \times 1$ vector and $\boldsymbol{\beta}_2$ is a $k_2 \times 1$ vector. The least squares estimator is given by:

$$\hat{\boldsymbol{\beta}} = \begin{bmatrix} \hat{\boldsymbol{\beta}}_1 \\ \hat{\boldsymbol{\beta}}_2 \end{bmatrix} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}.$$

We are interested in asking the question on whether can we retrieve the last k_2 entries of $\hat{\boldsymbol{\beta}}$ directly. The answer is yes we can by the **Frisch-Waugh-Lovell theorem**.

The first step is to regress $\mathbf{y}$ on $\mathbf{X}_1$ and retrieve the residuals. This is equivalent to applying the annihilator matrix

M_{X_1} onto y :

$$\hat{u}_{y.1} = \left[I - X_1(X_1^T X_1)^{-1} X_1^T \right] y = (I - P_{X_1})y = M_{X_1}y.$$

The second step is to regress X_2 on X_1 and retrieve the residuals:

$$\hat{u}_{2.1} = \left[I - X_1(X_1^T X_1)^{-1} X_1^T \right] X_2 = (I - P_{X_1})X_2 = M_{X_1}X_2.$$

The final step is to regress the residuals from the first step $\hat{u}_{y.1}$ on the residuals from the second step $\hat{u}_{2.1}$ and this gives us the OLS estimator:

$$\hat{\beta}_2 = (\hat{u}_{2.1}^T \hat{u}_{2.1})^{-1} \hat{u}_{2.1}^T \hat{u}_{y.1}.$$

This final equation is known as the **Frisch-Waugh-Lovell** theorem. □

We now look at a special case whereby we have a single explanatory variable (or two if you count the constant term on the intercept).

Theorem 381 (OLS with a Single Explanatory Variable)

Suppose we had the model:

$$y_i = \beta_1 + \beta_2 x_i + u_i \quad \text{for } i = 1, \dots, n.$$

Then, by the Frisch-Waugh-Lovell theorem,

$$\hat{\beta}_{FWL} = \frac{\sum_{i=1}^n (x_i - \bar{x}) y_i}{\sum_{i=1}^n (x_i - \bar{x})^2} \quad (307)$$

Proof. We are interested on how we can obtain an estimator of β_2 using a similar procedure as the Frisch-Waugh-Lovell theorem. We shall denote $\mathbf{1} = \begin{bmatrix} 1 & 1 & \dots & 1 \end{bmatrix}$ to be a $n \times 1$ matrix of 1's.

There are now four steps.

The first step is to regress y on the intercept term:

$$\begin{aligned} \hat{\beta}_{y.1} &= (\mathbf{1}^T \mathbf{1})^{-1} \mathbf{1}^T y \\ &= (n)^{-1} \mathbf{1}^T y \\ &= (n)^{-1} \sum_{i=1}^n y_i \\ &= \bar{y}. \end{aligned}$$

The second step is to regress X on the intercept term too:

$$\begin{aligned} \hat{\beta}_{X.1} &= (\mathbf{1}^T \mathbf{1})^{-1} \mathbf{1}^T X \\ &= (n)^{-1} \mathbf{1}^T X \\ &= (n)^{-1} \sum_{i=1}^n x_i \\ &= \bar{x}. \end{aligned}$$

The third step is to retrieve the residuals (observed minus fitted value) from step one and two so we have that the

individual observation is:

$$\begin{aligned}\hat{u}_{y \cdot 1, i} &= y_i - \bar{y} \cdot 1 = y_i - \bar{y} \\ \hat{u}_{x \cdot 1, i} &= x_i - \bar{x} \cdot 1 = x_i - \bar{x}\end{aligned}$$

We can stack these residuals into vectors. The final step is then to regress the first set of residuals $\hat{\mathbf{u}}_{y \cdot 1}$ on the other set of residuals $\hat{\mathbf{u}}_{x \cdot 1}$ to get:

$$\hat{\boldsymbol{\beta}}_{FWL} = (\hat{\mathbf{u}}_{x \cdot 1}^T \hat{\mathbf{u}}_{x \cdot 1})^{-1} \hat{\mathbf{u}}_{x \cdot 1}^T \hat{\mathbf{u}}_{y \cdot 1} = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2} =_{(1)} \frac{\sum_{i=1}^n (x_i - \bar{x})y_i}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

whereby $=_{(1)}$ holds by:

$$\begin{aligned}\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) &= \sum_{i=1}^n (x_i - \bar{x})y_i - \sum_{i=1}^n (x_i - \bar{x})\bar{y} \\ &= \sum_{i=1}^n (x_i - \bar{x})y_i - \bar{y} \sum_{i=1}^n (x_i - \bar{x}) \\ &= \sum_{i=1}^n (x_i - \bar{x})y_i - \bar{y}(n\bar{x} - n\bar{x}) \\ &= \sum_{i=1}^n (x_i - \bar{x})y_i - 0\end{aligned}$$

□

12.2.4 Correlation

Now, we revisit the concept of **correlation**. Recall that the population correlation is given by:

$$\text{Corr}[X, Y] = \frac{\text{Cov}[X, Y]}{\sqrt{\text{Var}[X]\text{Var}[Y]}}$$

We can now define the sample analogue.

Definition 382 (Sample Correlation)

The sample correlation between y_i and x_i is given by:

$$r_{y,x} = \text{Corr}(y_i, x_i) = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2 \sum_{i=1}^n (y_i - \bar{y})^2}} \quad (308)$$

We are now interested in **sample partial correlations**. Consider the partitioned equation again:

$$\mathbf{y} = \mathbf{X}_1 \boldsymbol{\beta}_1 + \mathbf{X}_2 \boldsymbol{\beta}_2 + \mathbf{u}$$

whereby $\mathbf{X}_2$ now has dimension $k_2 = 1$.

Definition 383 (Partial Correlation for one dimensional regressor)

The sample partial correlation is given by:

$$\text{Corr}(y_i, x_{2,i} | x_{1,i}) = \frac{\hat{u}_{2,1}^T \hat{u}_{y,1}}{\sqrt{(\hat{u}_{2,1}^T \hat{u}_{2,1})(\hat{u}_{y,1}^T \hat{u}_{y,1})}}. \quad (309)$$

whereby $\hat{u}_{2,1}$ is the residual vector from regressing $\mathbf{X}_2$ on $\mathbf{X}_1$ and $\hat{u}_{y,1}$ is the residual vector from regressing $\mathbf{y}$ on $\mathbf{X}_1$.

This definition of partial correlation follows a similar argument to the Frisch-Waugh-Lovell theorem whereby we are partialling out the effects of $\mathbf{X}_1$ on $\mathbf{X}_2$ and $\mathbf{y}$ respectively and looking at the correlation between them.

Remark 384. The sample correlation Eq. (308) is the partial correlation whereby we partial out the constant term $\mathbf{1}$ as $\hat{u}_{y,1} = y - \bar{y}$ and $\hat{u}_{x,1} = x - \bar{x}$,

$$r_{y,x} = \text{Corr}(y_i, x_i) = \text{Corr}(y_i, x_i | 1).$$

We can extend our definitions so far to looking at the correlation of $\mathbf{y}$ with a vector of explanatory variables $\mathbf{X}$.

Definition 385 (Sample Multiple Correlation)

Suppose that $\mathbf{x}$ is now a vector of explanatory variables. Then, the sample multiple correlation is:

$$\text{Corr}(\mathbf{y}, \mathbf{X}) = \left[\frac{\mathbf{y}^T \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}}{\mathbf{y}^T \mathbf{y}} \right]^{1/2} \quad (310)$$

Definition 386 (Sample Multiple Partial Correlations)

Suppose that $\mathbf{x}$ is now a vector of explanatory variables. Then, the sample multiple correlation is:

$$\text{Corr}(y_i, x_{2,i} | x_{1,i}) = \left[\frac{\hat{u}_{y,1}^T \hat{u}_{2,1} \left(\hat{u}_{2,1}^T \hat{u}_{2,1} \right)^{-1} \hat{u}_{2,1}^T \hat{u}_{y,1}}{\hat{u}_{y,1}^T \hat{u}_{y,1}} \right]^{1/2}. \quad (311)$$

Proposition 387 (Coefficient of Determination)

The coefficient of determination R^2 is the squared sample multiple partial correlation after partialling out the intercept. That is, for the model with an intercept:

$$\mathbf{y} = \beta_1 + \mathbf{X}\beta_2 + \mathbf{u}$$

we have that:

$$R^2 = \text{Corr}(y_i, x_i | 1)^2 = \frac{\hat{u}_{y,1}^T \hat{u}_{x,1} \left(\hat{u}_{x,1}^T \hat{u}_{x,1} \right)^{-1} \hat{u}_{x,1}^T \hat{u}_{y,1}}{\hat{u}_{y,1}^T \hat{u}_{y,1}} = \frac{ESS}{TSS}. \quad (312)$$

12.3 Autocorrelation

We move back to time series data. We want to capture the temporal dependence we observe in the data. To do so, we can modify the formula for sample correlation found in Eq. (308).

Definition 388 (Sample Autocorrelation)

For a fixed s , the sample autocorrelation is given by:

$$r_{0,s} = \frac{\sum_{t=s+1}^T (Y_t - \bar{Y}_{s+1}^T)(Y_{t-s} - \bar{Y}_1^{T-s})}{\sqrt{\sum_{t=s+1}^T (Y_t - \bar{Y}_{s+1}^T)^2 \sum_{t=s+1}^T (Y_{t-s} - \bar{Y}_1^{T-s})^2}} \quad (313)$$

whereby $\bar{Y}_a^b$ is the average of $Y_a, \dots, Y_b$.

Remark 389. The first bracket in the numerator looks the variation from time $s + 1$ up to time T . The second bracket in the numerator looks at the variation from time 1 to time $T - s$. Therefore, the first pair of observations is Y_{1+s} and Y_1 .

One can show that the sample autocorrelation is in fact asymptotically normal. However, we currently do not have all the techniques needed to show this.

Proposition 390 (Asymptotic Normality of Sample Autocorrelation under i.i.d)

Assume mean-zero i.i.d data and ignore mean correction. Then, the sample autocorrelation is given by:

$$r_{0,s} = \frac{\sum_{t=s+1}^T Y_t Y_{t-s}}{\sqrt{\sum_{t=s+1}^T Y_t^2 \sum_{t=s+1}^T Y_{t-s}^2}}.$$

Under proper rescaling, the sample autocorrelation is asymptotically normal:

$$\sqrt{T} r_{0,s} = \frac{T^{-1/2} \sum_{t=s+1}^T Y_t Y_{t-s}}{\sqrt{\left(T^{-1} \sum_{t=s+1}^T Y_t^2\right) \left(T^{-1} \sum_{t=s+1}^T Y_{t-s}^2\right)}} \xrightarrow{D} \mathcal{N}(0, 1).$$

Proof. (Sketch). We can apply a weak law of large numbers for the second moment in the denominator to both terms respectively. However, for the numerator, we have that the product $Y_t \cdot Y_{t-s}$ is a temporal sequence as $Y_t \cdot Y_{t-s}$ can share observations with other products in the sum. We will later establish that the scaled numerator converges to a normal distribution by a modified central limit theorem. Therefore, we can conclude from Slutsky's theorem that the scaled sample autocorrelation of an i.i.d sequence is asymptotically normal. $\square$

Remark 391. The issue in our proof is that we do not yet have a central limit theorem for dependent variables. We will work our way towards this in this course.

We now define our first stochastic process.

Definition 392 (Moving Average of Order 1)

A time series Y_t is a **moving average of order one** if:

$$Y_t = \epsilon_t + \theta \epsilon_{t-1} \quad (314)$$

for ϵ_t that are i.i.d with $\mathbb{E}[\epsilon_t] = 0$ and $\text{Var}[\epsilon_t] = \sigma^2$.

We state an important property that we will make use of repeatedly.

Theorem 393 (Property of Covariance)

The covariance operator has the property that for random variables X, Y, W, V :

$$\text{Cov}(aX + bY, cW + dV) = ac\text{Cov}(X, W) + ad\text{Cov}(X, V) + bc\text{Cov}(Y, W) + bd\text{Cov}(Y, V). \quad (315)$$

for constants a, b, c, d .

We can now state the variance and covariance of a MA(1) process.

Proposition 394 (Moments of a MA(1) Process)

Let Y_t be a MA(1) process. Then, we have that:

$$\mathbb{E}[Y_t] = 0 \quad (316)$$

$$\text{Var}(Y_t) = (1 + \theta^2)\sigma^2 \quad (317)$$

$$\text{Cov}(Y_t, Y_{t-1}) = \theta\sigma^2 \quad (318)$$

$$\text{Cov}(Y_t, Y_{t-j}) = 0 \quad \text{for } j \geq 2 \quad (319)$$

Proof. We repeatedly make use of the fact that ϵ_t are i.i.d.

$$\mathbb{E}[Y_t] = \mathbb{E}[\epsilon_t + \theta\epsilon_{t-1}] = 0.$$

$$\text{Var}[Y_t] = \mathbb{E}[Y_t^2] - \mathbb{E}[Y_t]^2 = \mathbb{E}[(\epsilon_t + \theta\epsilon_{t-1})^2] - 0 = \mathbb{E}[\epsilon_t^2] + 2\theta\mathbb{E}[\epsilon_t\epsilon_{t-1}] + \mathbb{E}[\theta^2\epsilon_{t-1}^2] = \sigma^2 + 0 + \theta^2\sigma^2.$$

$$\begin{aligned} \text{Cov}(Y_t, Y_{t-1}) &= \mathbb{E}[Y_t Y_{t-1}] \\ &= \mathbb{E}[(\epsilon_t + \theta\epsilon_{t-1})(\epsilon_{t-1} + \theta\epsilon_{t-2})] \\ &= \mathbb{E}[\epsilon_t\epsilon_{t-1}] + \theta\mathbb{E}[\epsilon_t\epsilon_{t-2}] + \theta\mathbb{E}[\epsilon_{t-1}^2] + \theta^2\mathbb{E}[\epsilon_{t-1}\epsilon_{t-2}] \\ &= 0 + \theta\mathbb{E}[\epsilon_{t-1}\epsilon_{t-1}] + 0 + 0 \\ &= \theta\text{Var}(\epsilon_{t-1}) \\ &= \theta\sigma^2 \end{aligned}$$

$$\text{Cov}(Y_t, Y_{t-j}) = \text{Cov}(\epsilon_t + \theta\epsilon_{t-1}, \epsilon_{t-j} + \theta\epsilon_{t-j-1}) = 0 \quad \text{for } j \geq 2$$

because the two sets of innovations are disjoint and independent when $j \geq 2$. □

Proposition 395 (Sample Autocorrelation for MA(1) Process)

The sample autocorrelation of a MA(1) process converges to the true population autocorrelation:

$$r_{0,1} \xrightarrow{P} \frac{\theta}{(1 + \theta^2)}. \quad (320)$$

Proof. (Sketch). We first compute the population autocorrelation:

$$\begin{aligned} \text{Corr}(Y_t, Y_{t-1}) &= \frac{\text{Cov}(Y_t, Y_{t-1})}{\text{Var}(Y_t)} \\ &= \frac{\theta \text{Var}(\epsilon_{t-1})}{(1 + \theta^2) \text{Var}(\epsilon_{t-1})} \\ &= \frac{\theta}{(1 + \theta^2)}. \end{aligned}$$

We will later argue using a modified law of large numbers that the sample autocorrelation converges to the true population autocorrelation. $\square$

Corollary 396

Higher-lag population autocorrelations for a MA(1) process are 0:

$$\text{Corr}(Y_t, Y_{t-j}) = 0 \quad \text{for } j \geq 2. \quad (321)$$

Proof. For $j \geq 2$, Y_t and Y_{t-j} depend on disjoint sets of independent innovations. Therefore their correlation is zero. $\square$

Definition 397 (Moving Average Of Order q)

A stochastic process Y_t is a moving average of order q if:

$$Y_t = \sum_{j=1}^q \theta_j \epsilon_{t-j} + \epsilon_t \quad (322)$$

for ϵ_t that are i.i.d.

Proposition 398 (Autocorrelation of MA(Q) Process)

Let Y_t be a moving average of order q. Then, the autocorrelation with all variables past the qth lag is equal to 0:

$$\text{Corr}(Y_t, Y_{t-s}) = 0$$

for $s \geq q + 1$.

12.4 Partial Autocorrelation

Now, suppose that we have our time series Y_t and this time, we want to find the **partial correlation** between Y_t and Y_{t-s} whilst partialling out all the Y_t 's in between Y_t and Y_{t-s} . These values will include $1, Y_{t-1}, Y_{t-2}, \dots, Y_{t-s+1}$.

Definition 399 (Partial Autocorrelation)

The partial autocorrelation between Y_t and Y_{t-s} is given by:

$$\text{Corr}(Y_t, Y_{t-s} | 1, Y_{t-1}, \dots, Y_{t-s+1}) = r_{0,s,1,\dots,s-1} = \frac{\hat{u}_{0|1,\dots,s-1}^T \hat{u}_{s|1,\dots,s-1}}{\sqrt{(\hat{u}_{0|1,\dots,s-1}^T \hat{u}_{0|1,\dots,s-1}) \cdot (\hat{u}_{s|1,\dots,s-1}^T \hat{u}_{s|1,\dots,s-1})}} \quad (323)$$

Remark 400. A partial correlogram simply partials out the effects of all lags in between our values of interest.

We can state useful properties of the partial autocorrelation.

Proposition 401 (Asymptotic Normality of Partial Autocorrelation)

Under the i.i.d white-noise null and for a fixed lag s , the partial autocorrelation is asymptotically normal.

$$T^{1/2}r_{0,s,1,\dots,s-1} \xrightarrow{D} \mathcal{N}(0, 1). \quad (324)$$

We now introduce another class of stochastic processes called autoregressive models.

Definition 402 (Autoregressive Model)

A time series Y_t is an autoregressive model of order 1 if:

$$Y_t = \alpha Y_{t-1} + \epsilon_t \quad (325)$$

for i.i.d ϵ_t .

Proposition 403 (Autocorrelation of AR(1) Process)

Let Y_t be a covariance-stationary AR(1) process with $|\alpha| < 1$. Then, the population autocorrelation is given by:

$$\text{Corr}(Y_t, Y_{t-s}) = \alpha^s \quad \text{for } s \geq 1. \quad (326)$$

Proof. First, we compute the population autocovariance for one lag $s = 1$:

$$\begin{aligned} \text{Cov}(Y_t, Y_{t-1}) &= \text{Cov}(\alpha Y_{t-1} + \epsilon_t, Y_{t-1}) \\ &= \text{Cov}(\alpha Y_{t-1}, Y_{t-1}) \\ &= \alpha \text{Cov}(Y_{t-1}, Y_{t-1}) \\ &= \alpha \text{Var}(Y_{t-1}) \end{aligned}$$

Then, we have that:

$$\text{Corr}(Y_t, Y_{t-1}) = \frac{\text{Cov}(Y_t, Y_{t-1})}{\text{Var}(Y_{t-1})} = \alpha.$$

By induction, this holds for:

$$\text{Corr}(Y_t, Y_{t-s}) = \alpha^s.$$

□

So for $|\alpha| < 1$ we have that the autocorrelation for an AR(1) process goes to 0 as $s \rightarrow \infty$.

Proposition 404 (Population Partial Autocorrelation of AR(1) Process)

The population partial autocorrelation of a covariance-stationary AR(1) process is 0:

$$\text{Corr}(Y_t, Y_{t-s} | Y_{t-1}, \dots, Y_{t-s+1}) = 0 \quad (327)$$

for $s \geq 2$.

Proof. Look at the AR(1) process:

$$Y_t = \alpha Y_{t-1} + \epsilon_t.$$

Y_t only depends on its past through Y_{t-1} . This means that:

$$f(Y_t|Y_{t-1}, Y_{t-2}) = f(Y_t|Y_{t-1}).$$

Then, we have that the conditional density is:

$$\begin{aligned} f(Y_t, Y_{t-2}|Y_{t-1}) &= \frac{f(Y_t, Y_{t-1}, Y_{t-2})}{f(Y_{t-1})} \\ &= \frac{f(Y_t, Y_{t-1}, Y_{t-2})}{f(Y_{t-1}, Y_{t-2})} \times \frac{f(Y_{t-1}, Y_{t-2})}{f(Y_{t-1})} \\ &= f(Y_t|Y_{t-1}, Y_{t-2}) \times f(Y_{t-2}|Y_{t-1}) \\ &= f(Y_t|Y_{t-1}) \times f(Y_{t-2}|Y_{t-1}) \end{aligned}$$

Hence, we have shown that:

$$f(Y_t, Y_{t-2}|Y_{t-1}) = f(Y_t|Y_{t-1}) \times f(Y_{t-2}|Y_{t-1})$$

and therefore we have shown that Y_t, Y_{t-2} are independent conditioned on Y_{t-1} . This implies that the partial autocorrelation is 0:

$$\text{Corr}(Y_t, Y_{t-s}|Y_{t-1}, \dots, Y_{t-s+1}) = 0.$$

□

Once we have partialled out the effects of Y_{t-1} , then we have that Y_t and Y_{t-2} are uncorrelated. This implies that the sample partial autocorrelation function of an AR(1) process should go to 0 by weak law of large numbers which we shall establish later.

To summarise the results for stochastic processes and correlations:

1. A cutoff in the ACF after q lags suggests an MA(q) series.
2. A cutoff in the PACF after p lags indicates AR(p) series.

12.5 Linear Projection in Time Series

Assumption 405 (Stationarity)

For this chapter, we assume that the time series Y_t is covariance stationary.

Cross Sectional Data Projection

First, we recall projections in cross-sectional data.

A linear predictor for Y is a function $X^T\beta$ for some $\beta \in \mathbb{R}^k$. The **mean squared prediction error** is:

$$S(\beta) = \mathbb{E}[(Y - X^T\beta)^2]. \quad (328)$$

The **best linear predictor** of Y given X is:

$$\mathcal{P}[Y|X] = X^T\beta$$

where β minimises the mean squared prediction error:

$$S(\beta) = \mathbb{E}[(Y - X^T\beta)^2].$$

Such a β is called the **linear projection coefficient**.

The mean squared prediction error is:

$$S(\beta) = \mathbb{E}[Y^2] - 2\beta^T \mathbb{E}[XY] + \beta^T \mathbb{E}[XX^T]\beta.$$

Take first order conditions with respect to β to get:

$$-2\mathbb{E}[XY] + 2\mathbb{E}[XX^T]\beta = 0.$$

We then solve to get:

$$\beta = (\mathbb{E}[XX^T])^{-1}\mathbb{E}[XY]. \quad (329)$$

Go back to our expression for best linear predictor and plug in the linear projection coefficient:

$$\mathcal{P}[Y|X] = X^T\beta = X^T(\mathbb{E}[XX^T])^{-1}\mathbb{E}[XY].$$

This is the **linear projection** of Y onto X .

The **projection error** is given by:

$$\epsilon = Y - \mathcal{P}[Y|X].$$

Time Series Data Projection

We now take this idea and think about projecting Y_t onto the closed linear span of its past:

$$\mathcal{H}_{t-1} = \overline{\text{span}}\{1, Y_{t-1}, Y_{t-2}, \dots\}.$$

We write the projection of Y_t onto $\mathcal{H}_{t-1}$ as:

$$\mathcal{P}[Y_t|\mathcal{H}_{t-1}]$$

This projection is unique and has a unique projection error (by the Hilbert projection theorem)

$$\epsilon_t = Y_t - \mathcal{P}[Y_t | \mathcal{H}_{t-1}].$$

$$\mathbb{E}[\epsilon_t] = 0$$

$$\mathbb{E}[\epsilon_t^2] = \sigma^2 < \infty$$

ϵ_t is serially uncorrelated.

If Y_t is strictly stationary, then $\mathcal{P}[Y_t | \mathcal{H}_{t-1}]$ and ϵ_t are strictly stationary too.

Theorem 406 (Covariance Stationary Process Have White Noise Projection Error)

If Y_t is a covariance-stationary process, then it has the projection equation:

$$Y_t = \mathcal{P}[Y_t | \mathcal{H}_{t-1}] + \epsilon_t \quad (330)$$

whereby the projection error ϵ_t satisfies:

$$\mathbb{E}[\epsilon_t] = 0$$

$$\mathbb{E}[\epsilon_{t-j}\epsilon_t] = 0 \quad j \geq 1$$

$$\mathbb{E}[\epsilon_t^2] = \sigma^2 < \infty.$$

If Y_t is strictly-stationary, then ϵ_t is strictly stationary.

Definition 407 (White Noise)

The process ϵ_t is a **white noise** if:

$$\mathbb{E}[\epsilon_t] = 0$$

$$\mathbb{E}[\epsilon_{t-j}\epsilon_t] = 0 \quad j \geq 1$$

$$\mathbb{E}[\epsilon_t^2] = \sigma^2 < \infty.$$

Proposition 408 (MDS is a White Noise)

A MDS with a finite, time-invariant second moment $\mathbb{E}[m_t^2] = \sigma^2$ is a white noise.

Remark 409. *The converse does not hold. A white noise is not necessarily a MDS.*

We list the order of processes from most restrictive to most general:

1. i.i.d sequence.
2. MDS.
3. White noise.

We have shown that covariance-stationary processes have a white noise projection error. We can extend upon this to an infinite number of projection errors.

Theorem 410 (Wold Decomposition)

Suppose that Y_t is a purely nondeterministic covariance-stationary process and $\mathbb{E}[\epsilon_t^2] = \sigma^2 < \infty$ where ϵ_t is the projection error. Then, the **Wold decomposition** of Y_t says that Y_t has the **linear representation**:

$$Y_t = \mu + \sum_{j=0}^{\infty} b_j \epsilon_{t-j} \quad (331)$$

whereby ϵ_{t-j} are the white noise projection errors:

$$\epsilon_{t-j} = Y_{t-j} - \mathcal{P}[Y_{t-j} | \mathcal{H}_{t-j-1}]$$

with

$$b_0 = 1, \quad \sum_{j=0}^{\infty} b_j^2 < \infty$$

and μ being the unconditional mean of Y_t .

Remark 411. For a general covariance-stationary process, the Wold decomposition may also contain a deterministic component given by the projection onto the remote past. The purely nondeterministic assumption rules out this additional component after demeaning.

The Wold decomposition says that a purely nondeterministic covariance-stationary process can be represented as a linear function of current and lagged projection errors. Here, μ is the **unconditional mean** of Y_t . This justifies the use of linear models as approximations to covariance-stationary processes.

We now define a useful algebraic construct for analysis in time series.

Definition 412 (Lag Operator)

The **lag operator** L satisfies:

$$L^k Y_t = Y_{t-k}. \quad (332)$$

We use the lag operator to compactly write the Wold representation.

Proposition 413 (Compact Form of Wold Representation)

The compact form of the Wold representation for a covariance stationary process is:

$$Y_t = \mu + b(L)\epsilon_t \quad (333)$$

where $b(L) = b_0 + b_1 L + b_2 L^2 + \dots$ is an infinite-order polynomial.

Proof. The Wold decomposition can then be written using the lag operator L as:

$$\begin{aligned} Y_t &= \mu + \sum_{j=0}^{\infty} b_j \epsilon_{t-j} \\ &= \mu + b_0 \epsilon_t + b_1 L \epsilon_t + b_2 L^2 \epsilon_t + \dots \\ &= \mu + (b_0 + b_1 L + b_2 L^2 + \dots) \epsilon_t \\ &= \mu + b(L) \epsilon_t \end{aligned}$$

where $b(z) = b_0 + b_1 z + b_2 z^2 + \dots$ is an infinite-order polynomial. □

To recap so far, first we have shown that Y_t satisfies a projection onto its infinite past. Then, the Wold representation showed that in fact, this projection is a linear function of lagged projection errors. Now, we want to show an alternative way to write Y_t involving lagged values of Y_t . We now write an infinite order representation of Y_t using lagged values of Y_t .

Theorem 414 (Autoregressive Wold Representation)

If Y_t is covariance-stationary with Wold representation $Y_t = b(L)\epsilon_t$ satisfying some technical conditions, then Y_t has the representation:

$$Y_t = \mu + \sum_{j=1}^{\infty} a_j Y_{t-j} + \epsilon_t \quad (334)$$

for some coefficients μ and a_j whereby

$$\sum_{k=0}^{\infty} k^s |a_k| < \infty$$

and whereby ϵ_{t-j} are white noise projection errors.

Remark 415. The technical assumptions says that for the Wold representation $Y_t = b(L)\epsilon_t$, the root of the polynomials $b(z)$ lie outside the unit circle $|z| > 1$. Furthermore, we require that the Wold coefficients satisfy:

$$\sum_{j=0}^{\infty} \left(\sum_{k=0}^{\infty} k^s b_{j+k} \right)^2 < \infty.$$

The summability of the Wold coefficients ensure convergence of the autoregressive coefficients a_j .

To summarise, we have shown two representations for a covariance stationary time series. The first is the Wold decomposition:

$$Y_t = \mu + \sum_{j=0}^{\infty} b_j \epsilon_{t-j}$$

whereby ϵ_{t-j} are white noise projection errors. The second representation is the autoregressive Wold representation:

$$Y_t = \mu + \sum_{j=1}^{\infty} a_j Y_{t-j} + \epsilon_t$$

whereby ϵ_t is also a white noise projection error. These two representations help us understand why we can approximate stationary time series with linear models, in particular autoregressive and moving-average models.

12.5.1 Linear Time Series Models

Definition 416 (AR(1) Process)

A process Y_t is an AR(1) process if:

$$Y_t = \alpha_0 + \alpha_1 Y_{t-1} + \epsilon_t \quad (335)$$

where ϵ_t is a strictly stationary white noise process.

If we back substitute Y_{t-1} , we end up with:

$$\begin{aligned} Y_t &= \alpha_0 + \alpha_1(\alpha_0 + \alpha_1 Y_{t-2} + \epsilon_{t-1}) + \epsilon_t \\ &= \alpha_0(1 + \alpha_1) + \alpha_1^2 Y_{t-2} + \alpha_1 \epsilon_{t-1} + \epsilon_t \end{aligned}$$

Continuously doing it for j -periods gives us:

$$Y_t = \alpha_0 \sum_{j=0}^{t-1} \alpha_1^j + \alpha_1^t Y_0 + \sum_{j=0}^{t-1} \alpha_1^j \epsilon_{t-j}.$$

Here, Y_t equals an intercept plus the scaled initial condition $\alpha_1^t Y_0$ and moving average expression. Suppose we can do this recursively into the infinite past and $|\alpha_1| < 1$. We then get the following result.

Proposition 417 (Convergent AR(1) Representation)

Suppose that Y_t is an AR(1) process and $|\alpha_1| < 1$. Then, its unique causal stationary solution has the convergent representation:

$$Y_t = \mu + \sum_{j=0}^{\infty} \alpha_1^j \epsilon_{t-j}. \quad (336)$$

where $\mu = \alpha_0 / (1 - \alpha_1)$. This solution is strictly stationary; a process started from an arbitrary fixed initial value converges to this stationary solution.

Proof. Look at the recursive form:

$$Y_t = \alpha_0 \sum_{j=0}^{t-1} \alpha_1^j + \alpha_1^t Y_0 + \sum_{j=0}^{t-1} \alpha_1^j \epsilon_{t-j}.$$

Use the result that if $|\beta| < 1$ then:

$$\sum_{k=0}^{\infty} \beta^k = \frac{1}{1 - \beta}.$$

Then since $|\alpha_1| < 1$

$$\begin{aligned} \alpha_0 \lim_{t \rightarrow \infty} \sum_{j=0}^{t-1} \alpha_1^j &= \frac{\alpha_0}{1 - \alpha_1} \equiv \mu. \\ \lim_{t \rightarrow \infty} \alpha_1^t Y_0 &= 0. \end{aligned}$$

□

We now use the lag operator to derive the same result. First, we can rewrite the AR(1) process as:

$$(1 - \alpha_1 L)Y_t = \alpha_0 + \epsilon_t.$$

Define the **autoregressive lag polynomial** $a(L) = 1 - \alpha_1 L$ to get:

$$a(L)Y_t = \alpha_0 + \epsilon_t.$$

As $a(L) = 1 - \alpha_1 L$, we want to know whether we can invert this. Recall that the geometric series is such that:

$$\sum_{j=0}^{\infty} x^j = \frac{1}{1 - x} = (1 - x)^{-1}$$

when $|x| < 1$. This gives us the following definition.

Definition 418 (Invertible)

A polynomial $\alpha(z)$ is invertible if

$$[\alpha(z)]^{-1} = \sum_{j=0}^{\infty} \alpha_j z^j \quad (337)$$

is absolutely convergent.

We now apply this to the AR(1) model. We currently have:

$$a(L)Y_t = \alpha_0 + \epsilon_t$$

where $a(L) = 1 - \alpha_1 L$. Now, if $a(L)$ is invertible, we have that:

$$[a(L)]^{-1} = [1 - \alpha_1 L]^{-1} = \frac{1}{1 - \alpha_1 L} = \sum_{j=0}^{\infty} \alpha_1^j L^j$$

Therefore, we apply this to the AR(1) model:

$$\begin{aligned} a(L)Y_t &= \alpha_0 + \epsilon_t \\ \Leftrightarrow Y_t &= a(L)^{-1}\alpha_0 + a(L)^{-1}\epsilon_t \\ \Leftrightarrow Y_t &= \left(\sum_{j=0}^{\infty} \alpha_1^j L^j \right) \alpha_0 + \left(\sum_{j=0}^{\infty} \alpha_1^j L^j \right) \epsilon_t \\ \Leftrightarrow Y_t &= \sum_{j=0}^{\infty} \alpha_1^j \alpha_0 + \left(\sum_{j=0}^{\infty} \alpha_1^j L^j \right) \epsilon_t \\ \Leftrightarrow Y_t &= \sum_{j=0}^{\infty} \alpha_1^j \alpha_0 + \sum_{j=0}^{\infty} \alpha_1^j \epsilon_{t-j} \\ \Leftrightarrow Y_t &= \frac{\alpha_0}{1 - \alpha_1} + \sum_{j=0}^{\infty} \alpha_1^j \epsilon_{t-j} \\ \Leftrightarrow Y_t &= \mu + \sum_{j=0}^{\infty} \alpha_1^j \epsilon_{t-j} \end{aligned}$$

We ended back at the convergent AR(1) representation!

We now move onto the class of MA(q) processes.

Definition 419 (MA(q) Process)

Y_t is a MA(q) process if it is of the form:

$$Y_t = \mu + \epsilon_t + \theta_1 \epsilon_{t-1} + \dots + \theta_q \epsilon_{t-q} \quad (338)$$

where ϵ_t is a strictly stationary white noise process.

Definition 420 (Infinite-Order Moving Average Process)

An infinite-order moving average process, denoted by $MA(\infty)$ is:

$$Y_t = \mu + \sum_{j=0}^{\infty} \theta_j \epsilon_{t-j}. \quad (339)$$

Remark 421. *This infinite-order moving average process is known as a **linear process**.*

Suppose we had the linear process:

$$Y_t = \sum_{j=0}^{\infty} \theta_j \epsilon_{t-j}.$$

The coefficients of the moving average representation are known as the impulse response function.

Definition 422 (Impulse Response Function)

The **impulse response function** is defined as the change in Y_{t+j} due to a shock in time:

$$\frac{\partial}{\partial \epsilon_t} Y_{t+j} = \theta_j. \quad (340)$$

θ_j is the magnitude of the impact of a time t shock on the time $t + j$ variable.

We have defined the IRF for a linear process. We now want to know how to calculate the impulse response b_j from the coefficients of an autoregressive model of order p :

$$Y_t = \alpha_1 Y_{t-1} + \dots + \alpha_p Y_{t-p} + \epsilon_t.$$

The impulse responses can be computed recursively from:

$$b_0 = 1, \quad b_j = \sum_{\ell=1}^{\min(p,j)} \alpha_{\ell} b_{j-\ell} \quad \text{for } j \geq 1. \quad (341)$$

12.6 Autoregression

We go back to our autoregressive equation of order 1.

Definition 423 (AR(1))

An AR(1) model is given by:

$$X_t = \alpha X_{t-1} + \epsilon_t \quad (342)$$

whereby $\epsilon_t \sim \mathcal{N}(0, \sigma^2)$ is i.i.d and independent of X_{t-j} for $j \geq 1$.

whereby X_0 is a **fixed value**. We can solve for a *solution* to this model.

Theorem 424 (Solution to AR(1))

The solution to an AR(1) model without deterministic terms is:

$$X_t = \sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0. \quad (343)$$

Proof. We can solve this equation through back-substitution:

$$\begin{aligned} X_t &= \alpha X_{t-1} + \epsilon_t \\ &= \alpha(\alpha X_{t-2} + \epsilon_{t-1}) + \epsilon_t \\ &= \alpha^2 X_{t-2} + \alpha \epsilon_{t-1} + \epsilon_t \\ &= \alpha^2(\alpha X_{t-3} + \epsilon_{t-2}) + \alpha \epsilon_{t-1} + \epsilon_t \\ &= \alpha^3 X_{t-3} + \alpha^2 \epsilon_{t-2} + \alpha \epsilon_{t-1} + \epsilon_t \end{aligned}$$

and via induction, we can arrive at:

$$X_t = \sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0.$$

□

Now, we recall a useful property of geometric series.

Lemma 425

The geometric series:

$$\sum_{j=0}^{t-1} \alpha^j = \begin{cases} \frac{1-\alpha^t}{1-\alpha} & \alpha \neq 1 \\ t & \alpha = 1 \end{cases} \quad (344)$$

$$\sum_{j=0}^{\infty} \alpha^j = \frac{1}{1-\alpha} \quad \text{for } |\alpha| < 1 \quad (345)$$

Theorem 426 (Conditional Mean and Variance of Conditional AR(1) Model)

Let X_t be an AR(1) model

$$X_t = \alpha X_{t-1} + \epsilon_t \quad (346)$$

with a fixed initial value X_0 . Suppose that ϵ_t is i.i.d and normally distributed $\mathcal{N}(0, \sigma^2)$. Assume that ϵ_t and the fixed initial value X_0 are independent. Finally, assume that $|\alpha| < 1$.

Then, the conditional mean is given by:

$$\mu_t \equiv \mathbb{E}[X_t|X_0] = \alpha^t X_0 \quad (347)$$

The conditional variance is given by:

$$\sigma_t^2 \equiv \text{Var}[X_t|X_0] = \sigma^2 \frac{1 - \alpha^{2t}}{1 - \alpha^2}. \quad (348)$$

Proof. First, recall our recursive solution:

$$X_t = \sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0.$$

We then take the conditional expectation of this:

$$\begin{aligned} \mathbb{E}[X_t|X_0] &= \mathbb{E}\left[\sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0 \middle| X_0\right] \\ &= \sum_{j=0}^{t-1} \alpha^j \mathbb{E}[\epsilon_{t-j}|X_0] + \alpha^t X_0 \\ &=_{(1)} \sum_{j=0}^{t-1} \alpha^j \mathbb{E}[\epsilon_{t-j}] + \alpha^t X_0 \\ &=_{(2)} \sum_{j=0}^{t-1} \alpha^j (0) + \alpha^t X_0 \\ &= \alpha^t X_0 \end{aligned}$$

whereby $=_{(1)}$ uses the independence of ϵ_t and X_0 and $=_{(2)}$ uses the mean of ϵ_t being 0.

For the conditional variance, we follow a similar approach by analysing the recursive solution.

$$\begin{aligned}
\text{Var}[X_t|X_0] &= \text{Var}\left[\sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0|X_0\right] \\
&= \text{Var}\left[\sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j}|X_0\right] \\
&= \sum_{j=0}^{t-1} \alpha^{2j} \text{Var}[\epsilon_{t-j}|X_0] \\
&= \sum_{j=0}^{t-1} \alpha^{2j} \text{Var}[\epsilon_{t-j}] \\
&= \sum_{j=0}^{t-1} \alpha^{2j} \sigma^2 \\
&= \sigma^2 \sum_{j=0}^{t-1} \alpha^{2j} \\
&= \sigma^2 \frac{1 - \alpha^{2t}}{1 - \alpha^2}.
\end{aligned}$$

□

We have derived the conditional mean and conditional variance for the AR(1) model. We now want to derive the limiting distribution now.

Theorem 427 (Asymptotic Normality of Conditional AR(1) Model)

Let X_t be an AR(1) model with a fixed initial value X_0 . Then, the conditional limiting distribution of X_t is:

$$(X_t|X_0) \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^2}{1 - \alpha^2}\right) \quad (349)$$

Proof. By the recursive solution, we know:

$$X_t = \sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0.$$

We also know that ϵ_t are i.i.d normal random variables. Therefore, the sum of i.i.d normal random variable is also a normal random variable. Furthermore, we have already derived the conditional mean and conditional variance for an AR(1) model, so we can conclude that:

$$X_t|X_0 \sim \mathcal{N}\left(\alpha^t X_0, \sigma^2 \frac{1 - \alpha^{2t}}{1 - \alpha^2}\right).$$

Now, we know that $|\alpha| < 1$ so we can take the limit to conclude that:

$$X_t|X_0 \rightarrow \mathcal{N}\left(0, \sigma^2 \frac{1}{1 - \alpha^2}\right).$$

□

Pausing to reflect what we have just done, we have shown that the AR(1) conditional model is asymptotically

normal:

$$(X_t|X_0) \xrightarrow{D} \mathcal{N}(0, \frac{\sigma^2}{1-\alpha^2})$$

The variance $\frac{\sigma^2}{1-\alpha^2}$ is the **stationary variance**. This is in contrast to the innovation variance σ^2 . The long-run variance relevant for the sample mean will instead be $\frac{\sigma^2}{(1-\alpha)^2}$.

We now want to examine the autocorrelation of the AR(1) model. Before we do that, we first need the autocovariance of the AR(1) model.

Proposition 428 (Conditional Autocovariance of AR(1) Model)

Let X_t be a conditional AR(1) process:

$$X_t = \alpha X_{t-1} + \epsilon_t$$

for a fixed X_0 , $|\alpha| < 1$ and $\epsilon_t \sim \mathcal{N}(0, \sigma^2)$ i.i.d and independent of all X_{t-j} for $j \geq 1$.

Then, for a fixed $s \geq 1$, the conditional autocovariance of the AR(1) model is:

$$\text{Cov}[X_t, X_{t-s}|X_0] = \alpha^s \text{Var}[X_{t-s}|X_0]. \quad (350)$$

Proof. We plug in the recursive formula for the AR(1) model into the conditional autocovariance function:

$$\begin{aligned} \text{Cov}[X_t, X_{t-s}|X_0] &= \text{Cov}\left[\sum_{j=0}^{s-1} \alpha^j \epsilon_{t-j} + \alpha^s X_{t-s}, X_{t-s}|X_0\right] \\ &=_{(1)} \text{Cov}[\alpha^s X_{t-s}, X_{t-s}|X_0] \\ &= \alpha^s \text{Cov}[X_{t-s}, X_{t-s}|X_0] \\ &= \alpha^s \text{Var}[X_{t-s}|X_0] \end{aligned}$$

where in $=_{(1)}$, we use the fact that ϵ_t was from time subscript $t - s + 1$ onward and hence is independent of X_{t-s} , implying a covariance of zero. $\square$

We now look at the autocorrelation of the AR(1) model whereby we perform asymptotics on the limiting distribution of the AR(1) model.

Proposition 429 (Autocorrelation for Limiting Distribution of AR(1) Model)

Let X_t be a conditional AR(1) process:

$$X_t = \alpha X_{t-1} + \epsilon_t$$

for a fixed X_0 , $|\alpha| < 1$ and $\epsilon_t \sim \mathcal{N}(0, \sigma^2)$ i.i.d and independent of all X_{t-j} for $j \geq 1$.

Then, for a fixed $s \geq 1$, the conditional autocorrelation for the limiting distribution of the AR(1) model is given by:

$$\text{Corr}(X_t, X_{t-s}|X_0) \rightarrow \alpha^s \quad (351)$$

as $t \rightarrow \infty$.

Proof. First, recall that the conditional autocovariance of the AR(1) model is:

$$\text{Cov}(X_t, X_{t-s}|X_0) = \alpha^s \text{Var}[X_{t-s}|X_0]$$

and the conditional variance of the AR(1) model is:

$$\text{Var}[X_{t-s}|X_0] = \sigma^2 \frac{1 - \alpha^{2(t-s)}}{1 - \alpha^2}.$$

Therefore, the limit of the conditional autocovariance is:

$$\lim_{t \rightarrow \infty} \alpha^s \text{Var}[X_{t-s}|X_0] = \lim_{t \rightarrow \infty} \alpha^s \sigma^2 \frac{1 - \alpha^{2(t-s)}}{1 - \alpha^2} = \alpha^s \frac{\sigma^2}{1 - \alpha^2}.$$

Finally, recall that we also have shown that the long-run variance was:

$$\lim_{t \rightarrow \infty} \text{Var}[X_t|X_0] = \frac{\sigma^2}{1 - \alpha^2}.$$

Now, we compute the formula for the conditional autocorrelation:

$$\begin{aligned} \text{Corr}(X_t, X_{t-s}|X_0) &= \frac{\text{Cov}(X_t, X_{t-s}|X_0)}{\sqrt{\text{Var}[X_t|X_0]\text{Var}[X_{t-s}|X_0]}} \\ &= \frac{\alpha^s \text{Var}[X_{t-s}|X_0]}{\sqrt{\text{Var}[X_t|X_0]\text{Var}[X_{t-s}|X_0]}} \\ &= \frac{\alpha^s \text{Var}[X_{t-s}|X_0]}{\sqrt{\text{Var}[X_t|X_0]\text{Var}[X_{t-s}|X_0]}} \end{aligned}$$

We then take limits of the numerator and denominator:

$$\frac{\lim_{t \rightarrow \infty} \alpha^s \text{Var}[X_{t-s}|X_0]}{\lim_{t \rightarrow \infty} \sqrt{\text{Var}[X_t|X_0]\text{Var}[X_{t-s}|X_0]}} = \frac{\alpha^s \frac{\sigma^2}{1 - \alpha^2}}{\frac{\sigma^2}{1 - \alpha^2}} = \alpha^s.$$

□

Hence, the autocorrelation for the limiting distribution is α^s .

Remark 430. When $0 < \alpha < 1$, we have an exponential decay in the autocorrelation of the AR(1) model.

12.7 Autoregressive Model With Intercepts and Time Trends

We now move onto analysing an AR(1) model with an intercept term.

Definition 431 (AR(1) Model with Intercept)

The AR(1) model with an intercept term is:

$$Y_t = \alpha Y_{t-1} + \epsilon_t + \mu \tag{352}$$

where μ is an intercept term and $\epsilon_t \sim \mathcal{N}(0, \sigma^2)$ i.i.d and independent of all Y_{t-j} for $j \geq 1$ and with a fixed initial value Y_0 .

The issue is that we cannot use our previous results as we now have a μ term that appears and by recursion, this μ term will also appear in X_{t-1} and so on. We now stress the following trick we will use to solve this issue because this trick is used all the time in time series. **The trick is to rewrite the AR(1) model with an intercept into a AR(1) model without an intercept plus the intercept term.** After we have rewritten the model, we can then apply the results we have already derived.

Proposition 432 (AR(1) Model without an Intercept)

Let Y_t be an AR(1) model with an intercept term. Then, we can rewrite Y_t as an AR(1) model without an intercept term by decomposing it into the sum of a mean-zero process X_t and long-run mean μ_Y :

$$X_t = \alpha X_{t-1} + \epsilon_t \quad (353)$$

$$Y_t = X_t + \mu_Y \quad (354)$$

$$\mu_Y = \frac{\mu}{1 - \alpha}. \quad (355)$$

Proof. First, let us take our AR(1) model and subtract off a constant μ_Y (this is an arbitrary constant for now but we can pin down what this constant will need to be later):

$$Y_t - \mu_Y = \alpha Y_{t-1} + \epsilon_t + \mu - \mu_Y$$

Now, we have a $Y_t - \mu_Y$ on the left hand side, and to get something similar on the right hand side (but dated one index back), we then add and subtract $\alpha\mu_Y$ on the right hand side:

$$Y_t - \mu_Y = \alpha Y_{t-1} + \epsilon_t + \mu - \mu_Y + \alpha\mu_Y - \alpha\mu_Y$$

We now group the terms as follows:

$$Y_t - \mu_Y = \alpha Y_{t-1} + \epsilon_t + \mu - \mu_Y - \alpha\mu_Y + \alpha\mu_Y$$

$$Y_t - \mu_Y = \alpha(Y_{t-1} - \mu_Y) + \epsilon_t + [\mu - (1 - \alpha)\mu_Y]$$

We can see that the terms in red actually form an AR(1) model without an intercept! To make this more explicit, we can define the variable:

$$X_t \equiv Y_t - \mu_Y$$

which then allows us to then rewrite the expression

$$Y_t - \mu_Y = \alpha(Y_{t-1} - \mu_Y) + \epsilon_t + [\mu - (1 - \alpha)\mu_Y]$$

to now be:

$$X_t = \alpha X_{t-1} + \epsilon_t + [\mu - (1 - \alpha)\mu_Y]$$

where the terms in blue are a constant that we now need to get rid off. That is, if we make the terms in blue = 0, then we end up with an AR(1) model without an intercept:

$$X_t = \alpha X_{t-1} + \epsilon_t$$

which is something we have already previously derived results for! Hence, going back to the blue term, we are required to pick a constant μ_Y such that:

$$\mu - (1 - \alpha)\mu_Y = 0$$

and clearly, this algebraically holds when:

$$\mu_Y = \frac{\mu}{1 - \alpha}.$$

Therefore, we set μ_Y to be this value. Hence, we can see that we have decomposed the AR(1) with an intercept

Y_t (just add and subtract μ_Y to Y_t) to be:

$$Y_t = \underbrace{(Y_t - \mu_Y)}_{=X_t} + \underbrace{\mu_Y}_{=\frac{\mu}{1-\alpha}} = X_t + \mu_Y$$

where X_t is an AR(1) without an intercept $X_t = \alpha X_{t-1} + \epsilon_t$ and μ_Y is just a constant. □

We can informally refer to this procedure as an AR(1) decomposition.

We now have the tools to derive the limiting distribution of an AR(1) model with an intercept term.

Proposition 433 (Limiting Distribution of AR(1) With Intercept)

The limiting distribution of an AR(1) model with an intercept Y_t is:

$$Y_t|Y_0 \xrightarrow{D} \mathcal{N}\left(\mu_Y, \frac{\sigma^2}{1-\alpha^2}\right) \quad (356)$$

Proof. Rewrite the AR(1) model with an intercept term:

$$Y_t = \alpha Y_{t-1} + \epsilon_t + \mu$$

into an AR(1) model without an intercept term plus a constant:

$$Y_t = X_t + \mu_Y$$

whereby $X_t = \alpha X_{t-1} + \epsilon_t$ and $\mu_Y = \frac{\mu}{1-\alpha}$. We have previously shown that the conditional distribution for an AR(1) without an intercept is:

$$X_t|X_0 \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^2}{1-\alpha^2}\right)$$

We can use this to derive the limiting distribution of Y_t as adding a constant μ_Y shifts the mean but does not change the variance and hence we can conclude:

$$Y_t|Y_0 \xrightarrow{D} \mathcal{N}\left(\mu_Y, \frac{\sigma^2}{1-\alpha^2}\right)$$

□

We now introduce the AR(1) model with an intercept **and** a linear time trend.

Definition 434 (AR(1) Model with Intercept and Linear Time Trend)

The AR(1) model with an intercept and linear time trend is given by:

$$Y_t = \alpha Y_{t-1} + \mu_c + \mu_I t + \epsilon_t \quad (357)$$

where μ_c is an intercept term, t is the linear time trend and $\epsilon_t \sim \mathcal{N}(0, \sigma^2)$ i.i.d and independent of all Y_{t-j} for $j \geq 1$ and with a fixed initial value Y_0 .

We now want to rewrite this model to remove the intercept and time trend.

Proposition 435 (AR(1) Model without an Intercept and Time Trend)

Let Y_t be an AR(1) model with an intercept term and a linear time trend term. Then, we can rewrite Y_t as an AR(1) model without an intercept term and linear time trend term by decomposing it into the sum of a mean-zero process X_t , long-run intercept mean τ_c and long-run time trend mean τ_l :

$$X_t = \alpha X_{t-1} + \epsilon_t \quad (358)$$

$$Y_t = X_t + \tau_c + \tau_l t \quad (359)$$

$$\tau_c = \frac{\mu_c - \alpha \tau_l}{1 - \alpha} \quad (360)$$

$$\tau_l = \frac{\mu_l}{1 - \alpha} \quad (361)$$

Proof. (Sketch). The proof is similar to removing the intercept from an AR(1) model. First, subtract τ_c and $\tau_l t$ from both sides of:

$$Y_t = \alpha Y_{t-1} + \mu_c + \mu_l t + \epsilon_t$$

to get:

$$Y_t - \tau_c - \tau_l t = \alpha Y_{t-1} + \mu_c + \mu_l t + \epsilon_t - \tau_c - \tau_l t$$

Now, add and subtract $\alpha \tau_c + \alpha \tau_l + \alpha \tau_l t$ to the right-hand side to get:

$$Y_t - \tau_c - \tau_l t = \alpha Y_{t-1} + \mu_c + \mu_l t + \epsilon_t - \tau_c - \tau_l t + \alpha \tau_c + \alpha \tau_l + \alpha \tau_l t - \alpha \tau_c - \alpha \tau_l - \alpha \tau_l t.$$

Now, start to group terms together to get:

$$Y_t - \tau_c - \tau_l t = \alpha(Y_{t-1} - \tau_c + \tau_l - \tau_l t) + \mu_c + \mu_l t - \tau_c - \tau_l t + \alpha \tau_l + \alpha \tau_l t - \alpha \tau_l + \epsilon_t.$$

and we end up with:

$$Y_t - \tau_c - \tau_l t = \alpha \left[Y_{t-1} - \tau_c - \tau_l(t-1) \right] + \left[\mu_c - \tau_c + \alpha \tau_c - \alpha \tau_l \right] + \left[\mu_l - \tau_l + \alpha \tau_l \right] t + \epsilon_t.$$

We then define the mean-zero process:

$$X_t \equiv Y_t - \tau_c - \tau_l t.$$

Therefore, we have that:

$$\underbrace{Y_t - \tau_c - \tau_l t}_{=X_t} = \alpha \underbrace{\left[Y_{t-1} - \tau_c - \tau_l(t-1) \right]}_{X_{t-1}} + \underbrace{\left[\mu_c - \tau_c + \alpha \tau_c - \alpha \tau_l \right]}_{=0} + \underbrace{\left[\mu_l - \tau_l + \alpha \tau_l \right]}_{=0} t + \epsilon_t.$$

to get an AR(1) process without an intercept or time-trend:

$$X_t = \alpha X_{t-1} + \epsilon_t$$

whereby we set the two terms equal to 0 to arrive at:

$$\mu_c - \tau_c + \alpha \tau_c - \alpha \tau_l = 0 \quad \leftrightarrow \quad \tau_c = \frac{\mu_c - \alpha \tau_l}{1 - \alpha}$$

and

$$\mu_l - \tau_l + \alpha \tau_l = 0 \quad \leftrightarrow \quad \tau_l = \frac{\mu_l}{1 - \alpha}.$$

□

We then can derive the limiting distribution of an AR(1) model with an intercept and time-trend.

Proposition 436 (Limiting Distribution of AR(1) With Intercept and Time-Trend)

The detrended AR(1) model has the conditional limiting distribution:

$$Y_t - (\tau_c + \tau_I t) | Y_0 \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^2}{1 - \alpha^2}\right) \quad (362)$$

whereby

$$\begin{aligned} \tau_c &= \frac{\mu_c - \alpha\tau_I}{1 - \alpha} \\ \tau_I &= \frac{\mu_I}{1 - \alpha}. \end{aligned}$$

Proof. First, decompose

$$Y_t = \alpha Y_{t-1} + \mu_c + \mu_I t + \epsilon_t$$

into an AR(1) without an intercept and time-trend:

$$X_t = \alpha X_{t-1} + \epsilon_t \quad (363)$$

$$Y_t = X_t + \tau_c + \tau_I t \quad (364)$$

$$\tau_c = \frac{\mu_c - \alpha\tau_I}{1 - \alpha} \quad (365)$$

$$\tau_I = \frac{\mu_I}{1 - \alpha} \quad (366)$$

Then, apply our result that the conditional limiting distribution of an AR(1) model is:

$$X_t | X_0 \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^2}{1 - \alpha^2}\right)$$

when $|\alpha| < 1$. Now if we look at:

$$Y_t = X_t + \tau_c + \tau_I t$$

we can then immediately conclude that the detrended process satisfies:

$$Y_t - (\tau_c + \tau_I t) | Y_0 \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^2}{1 - \alpha^2}\right).$$

□

Remark 437. In this scenario, we have that Y_t is **trend stationary**. That is, if we subtract a linear time trend from Y_t , we will get a stationary process.

Hence, the big takeaway in this section is how to analyse time series models with intercepts and time-trends by rewriting them in terms of models we already know properties about and adding a constant.

12.7.1 Random Walks

Definition 438 (Covariance Stationary)

A scalar process $\{y_t\}_{t \geq 1}$ is **covariance-stationary** if:

1. $\mathbb{E}[y_t] = \mu$
2. $\text{Cov}(y_t, y_{t+s}) = \gamma_s$ for all s .

Remark 439. *The important thing is that the first moment and covariance does not depend on time t .*

Definition 440 (White Noise)

The process u_t is a **white noise** if:

$$\mathbb{E}[u_t] = 0 \quad (367)$$

$$\mathbb{E}[u_{t-j}u_t] = 0 \quad j \geq 1 \quad (368)$$

$$\mathbb{E}[u_t^2] = \sigma_u^2 < \infty. \quad (369)$$

We now view what it means for a variable to be driven by a **stochastic trend**.

Definition 441 (AR(1) Process)

The AR(1) process is:

$$y_t = \alpha y_{t-1} + u_t \quad (370)$$

where u_t is a **white-noise process**.

When $\alpha = 1$, we get a random walk process.

Definition 442 (Random Walk)

A random walk process is:

$$y_t = y_{t-1} + u_t. \quad (371)$$

We can recursively substitute to get:

$$y_t = y_0 + \sum_{j=1}^t u_j.$$

Then:

$$\mathbb{E}[y_t] = \mathbb{E}[y_0]$$

$$\text{Var}[y_t] = t\sigma_u^2 + \text{Var}[y_0]$$

The variance of a random walk tends to infinity.

Proposition 443 (Persistence of Random Walk Process)

For any given integer s ,

$$\text{Corr}(y_t, y_{t+s}) \rightarrow 1 \quad (372)$$

as $t \rightarrow \infty$.

Proof. For simplicity, take y_0 to be fixed; this does not affect the limiting correlation. Then:

$$\text{Corr}(y_t, y_{t+s}) = \frac{\mathbb{E}\left[\left(\sum_{j=1}^t u_j\right)\left(\sum_{j=1}^{t+s} u_j\right)\right]}{\left[t\sigma_u^2(t+s)\sigma_u^2\right]^{1/2}} = \frac{t}{(t^2 + ts)^{1/2}} \rightarrow 1$$

as $t \rightarrow \infty$. □

Random variables y_t and y_{t+s} are still strongly correlated even when they are far apart in time. This indicates a strong persistence in the time process.

Definition 444 (Random Walk with Drift)

A random walk with a drift is:

$$y_t = \nu + y_{t-1} + u_t \quad (373)$$

where ν is a constant term and u_t is a white-noise process.

12.8 Stationarity

We now move onto the topic of stationarity. Essentially, stationarity imposes structure on our time series and makes things much easier for analysis.

Definition 445 (Strict Stationarity)

A process $\{X_t\}_t$ is strictly stationary if for any s , the joint distribution of $(X_t, X_{t+1}, \dots, X_{t+s})$ does not depend on time t .

Remark 446. Sometimes, an $*$ is used X_t^* to indicate that a process is a strictly stationary process.

Hence, if we have a strictly stationary process, then the joint distribution of $(X_{t+1}, \dots, X_{t+s})$ will be the same distribution of $(X_1, \dots, X_s)$. That is, if we have a block of variables and a joint density associated to this block, the distribution of the block will be the same no matter where we move it around.

Proposition 447 (Strict Stationarity is Preserved Under Measurability)

If X_t^* is a strictly stationary process and $g : \mathbb{R}^n \rightarrow \mathbb{R}$ is a measurable function, then

$$Z_t = g(X_t^*, X_{t-1}^*, \dots, X_{t-n}^*) \quad (374)$$

is a strictly stationary process.

We give an important class of strictly stationary processes.

Proposition 448 (i.i.d Sequences are Strictly Stationary)

If Y_t is i.i.d, then it is **strictly stationary**.

Example 449

Suppose we had i.i.d ϵ_t sequence. Then, ϵ_t is strictly stationary. Now, if we take a transformation $g(x) = x^2$ so that $g(\epsilon_t) = \epsilon_t^2$, we have that ϵ_t^2 is also a strictly stationary process.

We now define another notion of stationarity.

Definition 450 (Covariance Stationary)

A process $\{X_t\}_t$ is covariance stationary if for all s, t

1. $\mathbb{E}[X_t] = \mu$
2. $\text{Cov}(X_t, X_{t+s}) = \psi_s$.

Remark 451. *The first condition says that the expectation does not depend on the time t . Additionally, the covariance does not depend on time t as well. As a corollary, the unconditional variance does not depend on t , but the conditional variance can.*

Something important to note is that covariance stationarity is not preserved under transformations. An intuition for this is because if X_t has second moments, this does not imply that X_t^2 has second moments.

We now want to relate strict stationarity and covariance stationarity to each other. In fact, neither strict stationarity nor covariance stationarity is stronger than the other. However, if we impose some assumptions, we can relate the two concepts of stationarity.

Definition 452 (Gaussian Process)

A process X_t is a Gaussian process if for all s and time t , the joint distribution $(X_{t+1}, \dots, X_{t+s})$ is a Gaussian distribution.

We can now relate the two concepts of stationarity.

Proposition 453 (Sufficient Conditions for Strict Stationarity)

A Gaussian covariance stationary process is strictly stationary.

Proof. A finite-dimensional Gaussian distribution is characterised by its mean vector and covariance matrix. Hence, if these are invariant to a common shift in time, then the joint distribution is also invariant to that shift. Therefore, covariance stationarity implies strict stationarity for a Gaussian process. $\square$

Proposition 454 (Sufficient Conditions for Covariance Stationarity)

A strictly stationary process with second moments is a covariance stationary process.

Proof. We need to show that if X_t is strictly stationary, then its first moment and autocovariance is time-invariant.

First, we look at the first moment of X_t :

$$\mathbb{E}[X_t] = \int x f_t(x) dx \stackrel{=_{(1)}}{=} \int x f_1(x) dx = \mathbb{E}[X_1]$$

whereby $=_{(1)}$ uses strict stationarity of X_t . Hence, we have shown that:

$$\mathbb{E}[X_t] = \mathbb{E}[X_1]$$

for all time t .

Now look at the variance:

$$\begin{aligned}
\text{Var}[X_t] &= \mathbb{E}[X_t^2] - \mathbb{E}[X_t]^2 \\
&= \int x^2 f_t(x) dx - \mathbb{E}[X_t]^2 \\
&=_{(1)} \int x^2 f_1(x) dx - \mathbb{E}[X_t]^2 \\
&= \mathbb{E}[X_1^2] - \mathbb{E}[X_t]^2 \\
&= \text{Var}[X_1].
\end{aligned}$$

Hence, we have that $\text{Var}[X_t] = \text{Var}[X_1]$ for all time t .

Now, we look at the autocovariance:

$$\begin{aligned}
\text{Cov}[X_{t+h}, X_t] &= \mathbb{E}[X_{t+h}X_t] - \mathbb{E}[X_{t+h}]\mathbb{E}[X_t] \\
&= \int_{\mathbb{R}^2} xy f_{t+h,t}(x, y) dx dy - \mathbb{E}[X_{t+h}]\mathbb{E}[X_t] \\
&=_{(1)} \int_{\mathbb{R}^2} xy f_{1+h,1}(x, y) dx dy - \mathbb{E}[X_{t+h}]\mathbb{E}[X_t] \\
&= \mathbb{E}[X_{1+h}X_1] - \mathbb{E}[X_{t+h}]\mathbb{E}[X_t] \\
&=_{(2)} \mathbb{E}[X_{1+h}X_1] - \mathbb{E}[X_{1+h}]\mathbb{E}[X_1] \\
&= \text{Cov}[X_{1+h}, X_1]
\end{aligned}$$

whereby $=_{(1)}$ and $=_{(2)}$ uses strict stationarity. Hence we have shown that $\text{Cov}[X_{t+h}, X_t] = \text{Cov}[X_{1+h}, X_1]$ for all time t .

Therefore, we have shown that a strictly stationary process with second moments satisfy the conditions for covariance stationarity. $\square$

Example 455

(Cauchy Distribution). A series $\{X_t\}_{t \geq 1}$ can be strictly stationary if each X_t is Cauchy distributed. However, the Cauchy distribution does not have any moments. Therefore, such a sequence cannot be covariance stationary. That is why we require the assumption of the existence of second moments for a strictly stationary series to be covariance stationary.

Example 456

(MA(1) Process) Consider the MA(1) process:

$$X_t = \theta\epsilon_{t-1} + \epsilon_t$$

where $\{\epsilon_t\}_{t \geq 1}$ is i.i.d with $\mathbb{E}[\epsilon_t] = 0$ and $Var[\epsilon_t] = \sigma^2$.

Then, we can argue for covariance stationarity in two different ways.

The first way is to note that as ϵ_t is i.i.d, then we know that ϵ_t is strictly stationary. Furthermore, summing two strictly stationary series preserves strict stationarity so then:

$$X_t = \theta\epsilon_{t-1} + \epsilon_t$$

is strictly stationary. Finally, we know that X_t has finite second moments as ϵ_t as finite second moments and hence we can conclude that X_t is then covariance stationary.

The second way is verify the MA(1) process is covariance stationary directly. First, compute the first moment:

$$\mathbb{E}[X_t] = \mathbb{E}[\theta\epsilon_{t-1} + \epsilon_t] = 0.$$

The variance is:

$$Var[X_t] = \mathbb{E}[X_t^2] = (1 + \theta^2)\sigma^2.$$

Therefore, the covariance is:

$$Cov(X_t, X_{t-1}) = \theta\sigma^2$$

$$Cov(X_t, X_{t-s}) = 0 \quad \text{for } s \geq 2.$$

Hence we have shown directly that X_t is covariance stationary.

A lot of times, we have a partial sum of random variables:

$$\sum_{j=0}^t \alpha^j X_j$$

and we want to know whether can we simply take limits to get the series:

$$\sum_{j=0}^{\infty} \alpha^j X_j.$$

The following theorem tells us when we can do so.

Theorem 457 (Kolmogorov's 2 Series Theorem)

Let $\{Z_i\}_{i \geq 1}$ be independent random variables. Then suppose that:

1. the series $\sum_{i=0}^{\infty} \mathbb{E}[Z_i]$ converges
2. $\sum_{i=0}^{\infty} \text{Var}[Z_i] < \infty$.

Then, the series

$$\sum_{i=0}^{\infty} Z_i \quad (375)$$

converges almost surely.

Remark 458. We will give an example of applying this theorem in the next section.

12.9 Unconditional Stationary Autoregression

Previously, we have looked at the **conditional model** of AR(1) whereby:

$$X_t = \alpha X_{t-1} + \epsilon_t$$

where $|\alpha| < 1$ and we had a fixed initial value X_0 . This then gave us the limiting distribution:

$$X_t | X_0 \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^2}{1 - \alpha^2}\right).$$

We now discuss a second formulation of the AR(1) model known as the **unconditional** AR(1) model. The idea is that instead of fixing an initial value X_0 , we let X_0 be a random variable whose distribution is the stationary limiting distribution X_0^* .

Definition 459 (Unconditional AR(1) Model)

The **unconditional** AR(1) model is given by:

$$X_t^* = \alpha X_{t-1}^* + \epsilon_t \quad (376)$$

$$X_0^* \sim \mathcal{N}\left(0, \frac{\sigma^2}{1 - \alpha^2}\right) \quad (377)$$

whereby $\{X_t^*\}_{t \geq 1}$ is a stationary sequence, ϵ_t are i.i.d $\mathcal{N}(0, \sigma^2)$ and independent of X_0^* .

We can then derive the unconditional joint distribution and moments of such a process.

Proposition 460 (Joint Distribution of Unconditional AR(1) Model)

Let $\{X_t^*\}_{t \geq 1}$ be the unconditional AR(1) model. Then, for $s \geq 1$, the joint distribution is:

$$\begin{bmatrix} X_t^* \\ X_{t+s}^* \end{bmatrix} \sim \mathcal{N}\left(\begin{bmatrix} 0 \\ 0 \end{bmatrix}, \frac{\sigma^2}{1 - \alpha^2} \begin{bmatrix} 1 & \alpha^s \\ \alpha^s & 1 \end{bmatrix}\right) \quad (378)$$

Proof. First, we use backward substitution to get:

$$X_t^* = \sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0^*.$$

We now compute the **unconditional** moments of X_t^* . The first moment is:

$$\mathbb{E}[X_t^*] = \mathbb{E}\left[\sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0^*\right] = 0.$$

The variance is:

$$\begin{aligned} \text{Var}[X_t^*] &= \text{Var}\left[\sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0^*\right] \\ &=_{(1)} \text{Var}\left[\sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j}\right] + \text{Var}[\alpha^t X_0^*] \\ &= \sum_{j=0}^{t-1} \alpha^{2j} \sigma^2 + \alpha^{2t} \text{Var}[X_0^*] \\ &= \frac{\sigma^2(1 - \alpha^{2t})}{1 - \alpha^2} + \alpha^{2t} \frac{\sigma^2}{1 - \alpha^2} \\ &= \frac{\sigma^2}{1 - \alpha^2}. \end{aligned}$$

where $=_{(1)}$ uses the independence of ϵ_{t-j} and X_0 .

The autocovariance is:

$$\text{Cov}[X_t^*, X_{t+s}^*] = \mathbb{E}[X_t^* X_{t+s}^*] - \mathbb{E}[X_t^*] \mathbb{E}[X_{t+s}^*] = \mathbb{E}[X_t^* X_{t+s}^*] - 0$$

We now substitute in our recursive formula:

$$\begin{aligned} \text{Cov}[X_t^*, X_{t+s}^*] &= \mathbb{E}[X_t^* X_{t+s}^*] = \mathbb{E}\left[X_t^* \left(\sum_{j=0}^{s-1} \alpha^j \epsilon_{t+s-j} + \alpha^s X_t^*\right)\right] \\ &= \mathbb{E}\left[X_t^* \sum_{j=0}^{s-1} \alpha^j \epsilon_{t+s-j}\right] + \mathbb{E}[\alpha^s (X_t^*)^2] \\ &= \sum_{j=0}^{s-1} \alpha^j \mathbb{E}[X_t^* \epsilon_{t+s-j}] + \alpha^s \mathbb{E}[(X_t^*)^2] \\ &=_{(1)} 0 + \alpha^s \mathbb{E}[(X_t^*)^2] \\ &= \alpha^s \text{Var}[X_t^*] \\ &= \alpha^s \frac{\sigma^2}{1 - \alpha^2} \end{aligned}$$

where $=_{(1)}$ uses the independence of ϵ_{t+s-j} and X_t^* for $j = 0, \dots, s-1$. We see that the covariance does not depend on t .

Now, we look at the joint distribution of X_t^* and X_{t+s}^* :

$$X_t^* = \sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0^*$$

$$X_{t+s}^* = \sum_{j=0}^{t+s-1} \alpha^j \epsilon_{t+s-j} + \alpha^{t+s} X_0^*$$

We know that ϵ_t and X_0^* are normally distributed. Furthermore, we know that the sum of normal random variables is a normal random variable. Therefore, they are jointly normal with the joint distribution:

$$\begin{bmatrix} X_t^* \\ X_{t+s}^* \end{bmatrix} \sim \mathcal{N}\left(\begin{bmatrix} 0 \\ 0 \end{bmatrix}, \frac{\sigma^2}{1-\alpha^2} \begin{bmatrix} 1 & \alpha^s \\ \alpha^s & 1 \end{bmatrix}\right).$$

□

Remark 461. *Something important to notice is that this is an unconditional distribution. Previously, in the conditional AR(1) model, we derived a limiting conditional distribution. This is again due to the set up of the unconditional AR(1) model depending initially on a stationary X_0^* .*

Now, we know that the stationary solution to our unconditional AR(1) model $X_t = \alpha X_{t-1} + \epsilon_t$ is:

$$X_t = \sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0.$$

If $|\alpha| < 1$, we want to know whether can we take $t \rightarrow \infty$ and write:

$$X_t^* = \sum_{j=0}^{\infty} \alpha^j \epsilon_{t-j}.$$

The following results tells us yes.

Theorem 462 (Moving Average Representation of Unconditional AR(1) Model)

Let X_t be the unconditional AR(1) model. Then, if $|\alpha| < 1$, X_t has the moving average representation:

$$X_t^* = \sum_{j=0}^{\infty} \alpha^j \epsilon_{t-j}. \quad (379)$$

Proof. This follows from **Kolmogorov's 2 series theorem**. We verify that the conditions for the theorem hold. First, the innovations ϵ_{t-i} are independent by assumption. Now, we check the first condition:

$$\sum_{i=0}^{\infty} \mathbb{E}[\alpha^i \epsilon_{t-i}] = \sum_{i=0}^{\infty} \alpha^i \mathbb{E}[\epsilon_{t-i}] = 0$$

as $\mathbb{E}[\epsilon_i] = 0$.

We then check the variance:

$$\sum_{i=0}^{\infty} \text{Var}[\alpha^i \epsilon_{t-i}] = \sum_{i=0}^{\infty} \alpha^{2i} \sigma^2 = \frac{\sigma^2}{1-\alpha^2} < \infty. \quad (380)$$

Both conditions for Kolmogorov's 2 series theorem hold and hence:

$$\sum_{i=0}^{\infty} \alpha^i \epsilon_{t-i}$$

converges almost surely. □

To conclude this section, we look at the lag operator and show a different way to derive the moving average representation using lag operators.

Definition 463 (Lag Operator)

The lag operator is the operator L such that:

$$LX_t = X_{t-1}. \quad (381)$$

Now, if we have the AR(1) model:

$$X_t = \alpha X_{t-1} + \epsilon_t$$

we can express this in the **moving average representation**.

Theorem 464 (AR(1) Moving Average Representation)

Let X_t be an AR(1) model

$$X_t = \alpha X_{t-1} + \epsilon_t.$$

Then, if $|\alpha| < 1$, we can write X_t in its infinite order moving average representation:

$$X_t = \sum_{j=0}^{\infty} \alpha^j \epsilon_{t-j}. \quad (382)$$

Proof. First take the AR(1) model:

$$X_t = \alpha X_{t-1} + \epsilon_t.$$

We can rewrite this using the lag operator:

$$\begin{aligned} X_t - \alpha X_{t-1} &= \epsilon_t \\ \Leftrightarrow (1 - \alpha L)X_t &= \epsilon_t \end{aligned}$$

Now, let us define the polynomial:

$$A(z) = 1 - \alpha z \quad (383)$$

where if $|\alpha z| < 1$, then $A(z)$ has an inverse:

$$[A(z)]^{-1} = \sum_{j=0}^{\infty} \alpha^j z^j. \quad (384)$$

Then, we have that if we applied this to the AR(1) model:

$$(1 - \alpha L)X_t = \epsilon_t$$

$$\leftrightarrow X_t = [(1 - \alpha L)]^{-1} \epsilon_t$$

$$\leftrightarrow X_t = \sum_{j=0}^{\infty} \alpha^j \epsilon_{t-j}.$$

□

12.10 Asymptotic Theory for Time Series

First, we recall the definition of a σ – algebra.

Definition 465 (σ – algebra)

A σ – algebra is a collection of sets $\mathcal{F}$ such that:

1. $\Omega \in \mathcal{F}$
2. $A \in \mathcal{F}$ implies that $A^c \in \mathcal{F}$
3. If $A_1, A_2, \dots \in \mathcal{F}$ then

$$\bigcup_{i=1}^{\infty} A_i \in \mathcal{F}.$$

Definition 466 (Filtration)

A filtration is a sequence of σ – algebra's $\mathcal{F}_t$ such that:

$$\mathcal{F}_t \subseteq \mathcal{F}_{t+1} \tag{385}$$

for all time.

A filtration captures the idea that the amount of information you have is *growing* over time.

Definition 467 (Martingale)

A process $\{X_t\}_{t \geq 1}$ is a martingale if

$$\mathbb{E}[X_{t+1}|\mathcal{F}_t] = X_t. \tag{386}$$

Remark 468. *The best guess for the stock price tomorrow is the stock price today, on average.*

Now, suppose we construct a new process consisting of taking the difference of martingales:

$$m_t = X_t - X_{t-1}$$

whereby $\{X_t\}_{t \geq 1}$ is a martingale. Then, the conditional expectation of m_t is:

$$\begin{aligned} \mathbb{E}[m_t|\mathcal{F}_{t-1}] &= \mathbb{E}[X_t - X_{t-1}|\mathcal{F}_{t-1}] \\ &= \mathbb{E}[X_t|\mathcal{F}_{t-1}] - \mathbb{E}[X_{t-1}|\mathcal{F}_{t-1}] \\ &= \mathbb{E}[X_t|\mathcal{F}_{t-1}] - X_{t-1} \\ &= X_{t-1} - X_{t-1} = 0. \end{aligned}$$

We now define an important concept that we will use frequently.

Definition 469 (Martingale Difference Sequence)

A pair $(m_t, \mathcal{F}_t)$ for $t \in \mathbb{N}$ is a **martingale difference sequence** if:

1. m_t is $\mathcal{F}_t$ -measurable
2. $\mathbb{E}[|m_t|] < \infty$
3. $\mathbb{E}[m_t | \mathcal{F}_{t-1}] = 0$.

As we will soon see, martingale differences are a relaxation of the assumption of independence. To see this, we derive properties of the MDS.

Proposition 470 (Unconditional Mean of MDS)

Let $(m_t, \mathcal{F}_t)_{t \geq 1}$ be a MDS. The unconditional mean of a MDS is:

$$\mathbb{E}[m_t] = 0. \quad (387)$$

Proof. Apply the law of iterated expectations and use the property of a MDS:

$$\mathbb{E}[m_t] = \mathbb{E}[\mathbb{E}[m_t | \mathcal{F}_{t-1}]] = \mathbb{E}[0] = 0.$$

□

Proposition 471 (Conditional Mean of a MDS)

Let $(m_t, \mathcal{F}_t)_{t \geq 1}$ be a MDS. Suppose that $s < t$. Then, the conditional mean is:

$$\mathbb{E}[m_t | \mathcal{F}_s] = 0. \quad (388)$$

Proof. Apply the law of iterated expectations:

$$\begin{aligned} \mathbb{E}[m_t | \mathcal{F}_s] &= \mathbb{E}[\mathbb{E}[m_t | \mathcal{F}_{t-1}] | \mathcal{F}_s] \\ &= \mathbb{E}[0 | \mathcal{F}_s] \\ &= 0 \end{aligned}$$

□

Proposition 472 (Autocovariance of a MDS)

Let $(m_t, \mathcal{F}_t)_{t \geq 1}$ be a MDS. Assume that $\mathbb{E}[m_t^2] < \infty$. Then, for $s < t$,

$$\text{Cov}[m_s, m_t] = 0. \quad (389)$$

Proof. Use the formula for covariance and the fact that the conditional mean of a MDS is 0:

$$\begin{aligned}
\text{Cov}[m_s, m_t] &= \mathbb{E}[m_s m_t] - \mathbb{E}[m_s] \mathbb{E}[m_t] \\
&= \mathbb{E}[m_s m_t] - 0 \\
&= \mathbb{E}[\mathbb{E}[m_s m_t | \mathcal{F}_s]] \\
&= \mathbb{E}[m_s \underbrace{\mathbb{E}[m_t | \mathcal{F}_s]}_{=0}] \\
&= \mathbb{E}[m_s \cdot 0] \\
&= 0
\end{aligned}$$

□

Proposition 473 (A MDS Process is Serially Uncorrelated)

A MDS $(m_t, \mathcal{F}_t)_{t \geq 1}$ with $\mathbb{E}[m_t^2] < \infty$ is a serially **uncorrelated** sequence.

Proof. We have shown that the covariance of a MDS is zero. Therefore it implies it is uncorrelated. □

Hence, we can see that with a MDS, it is a relaxation of an i.i.d sequence.

Remark 474. Note that the converse does not hold. That is, a serially uncorrelated process is not necessarily a MDS.

Proposition 475 (i.i.d Sequence with Mean Zero)

Every i.i.d sequence with mean zero is a MDS.

Proof. We use the fact that m_t is independent and has mean zero to see:

$$\mathbb{E}[m_t | \mathcal{F}_{t-1}] = \mathbb{E}[m_t] = 0.$$

□

12.11 Limit Theorems for Martingale Difference Sequences

We can now state limit theorems for MDS.

Theorem 476 (Law of Large Numbers for Martingale Difference Sequence)

Let $(m_t, \mathcal{F}_t)$ be a martingale difference sequence. If either of:

1. $\sum_{t=1}^{\infty} \frac{\mathbb{E}[|m_t|^{1+p}]}{t^{1+p}} < \infty$ for some $p \in [0, 1]$
2. $\lim_{T \rightarrow \infty} \max_{t \leq T} [\mathbb{E}[|m_t|^{1+p}]] < \infty$ for some $p \in [0, 1]$,

then we have that:

$$\frac{1}{T} \sum_{t=1}^T m_t \xrightarrow{P} \mathbb{E}[m_t] = 0. \quad (390)$$

Remark 477. Recall that we have shown that the unconditional mean of a MDS $\mathbb{E}[m_t] = 0$. Hence, this result makes sense.

We now give an example of how to check these conditions.

Example 478

Let m_t be an i.i.d sequence with $\mathbb{E}[m_t^2] = \sigma^2$. Then, let $p = 1$. We then check the first condition.

$$\begin{aligned} \sum_{t=1}^{\infty} \frac{\mathbb{E}[|m_t|^2]}{t^2} &= \sum_{t=1}^{\infty} \frac{\mathbb{E}[m_t^2]}{t^2} \\ &= \sum_{t=1}^{\infty} \frac{\sigma^2}{t^2} \\ &= \sigma^2 \sum_{t=1}^{\infty} \frac{1}{t^2} \\ &= \sigma^2 \frac{\pi^2}{6} < \infty. \end{aligned}$$

Hence, the first condition is satisfied so we can apply the LLN to m_t .

We now show an application of the second condition.

$$\begin{aligned} \lim_{T \rightarrow \infty} \max_{t \leq T} [\mathbb{E}[|m_t|^2]] &= \lim_{T \rightarrow \infty} \max_{t \leq T} \sigma^2 \\ &= \sigma^2 < \infty. \end{aligned}$$

Therefore, we see that the second condition is satisfied. Therefore we can apply the law of large number to m_t .

We can now state the central limit theorem for MDS.

Theorem 479 (Central Limit Theorem for Martingale Difference Sequence)

Let $(m_t, \mathcal{F}_t)$ be a martingale difference sequence. Let $S_T^2 = \sum_{t=1}^T \mathbb{E}[m_t^2]$. Then, if:

1. $\frac{1}{S_T^2} \sum_{t=1}^T \mathbb{E}[m_t^2 | \mathcal{F}_{t-1}] \xrightarrow{P} 1$ and
2. $\frac{1}{S_T^2} \sum_{t=1}^T \mathbb{E}[m_t^2 1\{|m_t| > \delta S_T\}] \rightarrow 0$ for every $\delta > 0$,

then we have that:

$$\frac{1}{S_T} \sum_{t=1}^T m_t \xrightarrow{D} \mathcal{N}(0, 1). \quad (391)$$

Remark 480. The second condition is known as a **Lindeberg condition**. Its intuition is that no unusually large martingale difference should dominate the normalized sum.

We give an example for checking this.

Example 481

Suppose that m_t is i.i.d with mean zero and constant second moment $\mathbb{E}[m_t^2] = \sigma^2$. Then,

$$S_T^2 = \sum_{t=1}^T \mathbb{E}[m_t^2] = \sum_{t=1}^T \sigma^2 = T\sigma^2.$$

Then,

$$S_T = \sqrt{T\sigma^2}$$

and hence:

$$\frac{1}{\sqrt{T\sigma^2}} \sum_{t=1}^T m_t \xrightarrow{D} \mathcal{N}(0, 1).$$

12.12 Limit Theorems for AR(1) Model

We have stated the limit theorems for MDS. We now want to derive limit theorems for the AR(1) model which is what is generally used in this course. To derive these AR(1) limit theorem, we can use the limit theorems for MDS.

First we give a motivation for why we need limit theorems for AR(1) models. Suppose we are interested in testing $H_0 : \mu_Y = 0$ in the model:

$$Y_t = \mu_Y + u_t$$

where we have a correlated error term:

$$u_t = \alpha u_{t-1} + \epsilon_t.$$

If we **incorrectly assume** that u_t is i.i.d, then we would then incorrectly assume that Y_t is i.i.d. Then, when testing $H_0 : \mu_Y = 0$, we would construct the ordinary test statistic:

$$t = \frac{T^{1/2}(\bar{Y} - \mu_Y)}{S_Y} \xrightarrow{D} \mathcal{N}(0, 1)$$

where S_Y is the standard deviation of Y . However, as we will later prove, if we assumed instead correlated error terms, we would instead end with the limiting distribution:

$$t = \frac{T^{1/2}(\bar{Y} - \mu_Y)}{S_Y} \xrightarrow{D} \mathcal{N}(0, \frac{1+\alpha}{1-\alpha}).$$

We can see that we end up with two different distributions:

$$\mathcal{N}(0, 1) \quad \text{and} \quad \mathcal{N}(0, \frac{1+\alpha}{1-\alpha}).$$

We plot these distributions in Fig. 2. If $\alpha > 0$, then $\frac{1+\alpha}{1-\alpha} > 1$, so the correctly derived limiting distribution has thicker tails than $\mathcal{N}(0, 1)$. Using standard-normal critical values then leads to rejecting the null hypothesis too often. If $\alpha < 0$, the variance ratio is below one and the standard-normal test is instead conservative. The two distributions coincide when $\alpha = 0$, in which case the errors are uncorrelated.

With this motivation in mind, we now aim to derive the limit theorems for AR(1) models.

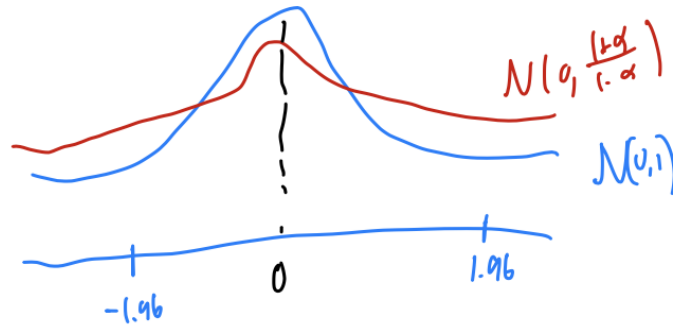


Figure 2: Comparison of the two limiting distributions of the same test statistic under different assumptions.

Theorem 482 (Law of Large Numbers for AR(1))

Let X_t be an AR(1) process

$$X_t = \alpha X_{t-1} + \epsilon_t.$$

Assume that

1. $|\alpha| < 1$
2. ϵ_t is i.i.d with mean zero, variance σ^2 , and finite fourth moment
3. X_0 is fixed, or is independent of the innovations and has a finite fourth moment.

Then, for any fixed non-negative integer j , the three laws of large numbers for AR(1) processes are:

$$\frac{1}{T} \sum_{t=1}^T X_{t-j} \xrightarrow{P} 0 \quad (392)$$

$$\frac{1}{T} \sum_{t=1}^T X_{t-j}^2 \xrightarrow{P} \frac{\sigma^2}{1-\alpha^2} \equiv \sigma_X^2 \quad (393)$$

$$\frac{1}{T} \sum_{t=1}^T X_{t-1} \epsilon_t \xrightarrow{P} \mathbb{E}[X_{t-1} \epsilon_t] = 0. \quad (394)$$

Proof. The general strategy for proving each of these results is to first rewrite the expression as a MDS. Then, apply either the MDS law of large numbers or MDS central limit theorem.

Proof of 1

First subtract αX_t from both sides of the AR(1) model:

$$\begin{aligned} X_t &= \alpha X_{t-1} + \epsilon_t \\ \Leftrightarrow (1-\alpha)X_t &= -\alpha(X_t - X_{t-1}) + \epsilon_t \\ \Leftrightarrow X_t &= \frac{-\alpha}{(1-\alpha)}(X_t - X_{t-1}) + \frac{1}{(1-\alpha)}\epsilon_t \\ \Leftrightarrow X_t &= \frac{\alpha}{(\alpha-1)}(X_t - X_{t-1}) + \frac{1}{(1-\alpha)}\epsilon_t \end{aligned}$$

where the second last equation is valid as $|\alpha| < 1$. Then we sum both sides up to T :

$$\begin{aligned}\sum_{t=1}^T X_t &= \frac{\alpha}{\alpha-1} \sum_{t=1}^T (X_t - X_{t-1}) + \frac{1}{1-\alpha} \sum_{t=1}^T \epsilon_t \\ \Leftrightarrow \sum_{t=1}^T X_t &= \frac{\alpha}{\alpha-1} (X_T - X_0) + \frac{1}{1-\alpha} \sum_{t=1}^T \epsilon_t.\end{aligned}$$

where we used the fact that $\sum_{t=1}^T (X_t - X_{t-1})$ is a telescoping sum. Then divide by T on both sides:

$$\frac{1}{T} \sum_{t=1}^T X_t = \frac{1}{T} \frac{\alpha}{\alpha-1} (X_T - X_0) + \frac{1}{1-\alpha} \frac{1}{T} \sum_{t=1}^T \epsilon_t.$$

The final step is to now show that in:

$$\frac{1}{T} \sum_{t=1}^T X_t = \underbrace{\frac{1}{T} \frac{\alpha}{\alpha-1} (X_T - X_0)}_{(*)} + \underbrace{\frac{1}{1-\alpha} \frac{1}{T} \sum_{t=1}^T \epsilon_t}_{(**)}.$$

that $(*) = o_p(1)$ and $(**) \xrightarrow{P} 0$.

First, we look at $(*)$:

$$\frac{\alpha}{\alpha-1} \frac{1}{T} (X_T - X_0)$$

As X_0 does not depend on T , it is a fixed random variable and therefore $X_0 = O_p(1)$ and therefore $\frac{1}{T} X_0 = o_p(1)$. We now look at the second term:

$$\frac{1}{T} X_T.$$

We need to show that this is also $o_p(1)$ which means showing:

$$\lim_{T \rightarrow \infty} \mathbb{P}\left[\frac{|X_T|}{T} > \delta\right] = 0.$$

To do so, we can apply **Markov's inequality**:

$$\begin{aligned}\mathbb{P}\left[\frac{|X_T|}{T} > \delta\right] &\leq \frac{1}{\delta^2 T^2} \mathbb{E}[X_T^2] \\ &\leq \frac{C}{\delta^2 T^2} \rightarrow 0\end{aligned}$$

as $T \rightarrow \infty$, where stability and the moment assumptions imply $\sup_T \mathbb{E}[X_T^2] \leq C < \infty$. Hence, we have just shown that $\frac{1}{T} X_T = o_p(1)$. We then use the fact that:

$$o_p(1) + o_p(1) = o_p(1)$$

to conclude that

$$\frac{\alpha}{\alpha-1} \frac{1}{T} (X_T - X_0) = o_p(1).$$

Now look at $(**)$:

$$\frac{1}{T} \sum_{t=1}^T \epsilon_t.$$

This is a MDS. We can therefore apply the **MDS law of large numbers**:

$$\frac{1}{T} \sum_{t=1}^T \epsilon_t \xrightarrow{P} \mathbb{E}[\epsilon_t] = 0.$$

Hence, we shown the desired results for (*) and (**) so we have that:

$$\begin{aligned} \frac{1}{T} \sum_{t=1}^T X_t &= \underbrace{\frac{1}{T} \frac{\alpha}{\alpha-1} (X_T - X_0)}_{(*)} + \underbrace{\frac{1}{1-\alpha} \frac{1}{T} \sum_{t=1}^T \epsilon_t}_{(**)} \\ &\leftrightarrow \frac{1}{T} \sum_{t=1}^T X_t = o_p(1) + o_p(1) \end{aligned}$$

and hence:

$$\frac{1}{T} \sum_{t=1}^T X_t \xrightarrow{P} 0.$$

Proof of 3

We first prove the third AR(1) law of large numbers before proving the second AR(1) law of large numbers. We want to show that:

$$\frac{1}{T} \sum_{t=1}^T X_{t-1} \epsilon_t \xrightarrow{P} \mathbb{E}[X_{t-1} \epsilon_t] = 0.$$

First, we notice that the summand is a martingale difference sequence:

$$\mathbb{E}[X_{t-1} \epsilon_t | \mathcal{F}_{t-1}] = X_{t-1} \mathbb{E}[\epsilon_t | \mathcal{F}_{t-1}] = 0.$$

Therefore, we want to apply the law of large numbers for martingale differences to $X_{t-1} \epsilon_t$. We check the bounded-moment condition with $p = 1$:

$$\begin{aligned} \mathbb{E}[(X_{t-1} \epsilon_t)^2] &= \mathbb{E}[\mathbb{E}[(X_{t-1} \epsilon_t)^2 | \mathcal{F}_{t-1}]] \\ &= \mathbb{E}[X_{t-1}^2 \mathbb{E}[\epsilon_t^2 | \mathcal{F}_{t-1}]] \\ &= \sigma^2 \mathbb{E}[X_{t-1}^2]. \end{aligned}$$

Stability and the moment assumptions give $\sup_t \mathbb{E}[X_{t-1}^2] < \infty$, so

$$\sup_t \mathbb{E}[(X_{t-1} \epsilon_t)^2] = \sigma^2 \sup_t \mathbb{E}[X_{t-1}^2] < \infty.$$

Therefore, we can apply the WLLN for MDS and conclude that:

$$\frac{1}{T} \sum_{t=1}^T X_{t-1} \epsilon_t \xrightarrow{P} \mathbb{E}[X_{t-1} \epsilon_t] = 0.$$

Proof of 2

We now want to show:

$$\frac{1}{T} \sum_{t=1}^T X_{t-j}^2 \xrightarrow{P} \mathbb{E}[X_t^2] = \text{Var}[X_t] = \frac{\sigma^2}{1-\alpha^2} \equiv \sigma_X^2$$

First, if we expand out the square we get:

$$X_t^2 = \alpha^2 X_{t-1}^2 + \epsilon_t^2 + 2\alpha X_{t-1}\epsilon_t.$$

Now, if we take the sum of both sides, we get:

$$\sum_{t=1}^T X_t^2 = \alpha^2 \sum_{t=1}^T X_{t-1}^2 + \sum_{t=1}^T \epsilon_t^2 + 2\alpha \sum_{t=1}^T X_{t-1}\epsilon_t.$$

We then factorise this to get:

$$\sum_{t=1}^T X_t^2 = \alpha^2 \left[\sum_{t=1}^T X_t^2 + X_0^2 - X_T^2 \right] + \sum_{t=1}^T \epsilon_t^2 + 2\alpha \sum_{t=1}^T X_{t-1}\epsilon_t.$$

whereby we used the fact that: $\sum_{t=1}^T X_{t-1}^2 = \sum_{t=1}^T X_t^2 + X_0^2 - X_T^2$. Re-arrange to get:

$$\sum_{t=1}^T X_t^2 - \alpha^2 \sum_{t=1}^T X_t^2 = \alpha^2 \left[X_0^2 - X_T^2 \right] + \sum_{t=1}^T \epsilon_t^2 + 2\alpha \sum_{t=1}^T X_{t-1}\epsilon_t$$

Factorise and divide through to get:

$$\sum_{t=1}^T X_t^2 = \frac{\alpha^2}{1-\alpha^2} \left[X_0^2 - X_T^2 \right] + \frac{\sum_{t=1}^T \epsilon_t^2}{1-\alpha^2} + \frac{2\alpha}{1-\alpha^2} \sum_{t=1}^T X_{t-1}\epsilon_t$$

We then divide by T on both sides:

$$\frac{1}{T} \sum_{t=1}^T X_t^2 = \frac{\alpha^2}{1-\alpha^2} \left[\frac{X_0^2}{T} - \frac{X_T^2}{T} \right] + \frac{1}{T} \frac{\sum_{t=1}^T \epsilon_t^2}{1-\alpha^2} + \frac{1}{T} \frac{2\alpha}{1-\alpha^2} \sum_{t=1}^T X_{t-1}\epsilon_t$$

We now look at each term respectively. First, we know that:

$$\frac{X_0^2}{T} = o_p(1)$$

since $o(1)O_p(1) = o(1) = o_p(1)$. The second term is:

$$\frac{X_T^2}{T} = o_p(1)$$

through Markov's inequality. We then can conclude that:

$$\left[\frac{X_0^2}{T} - \frac{X_T^2}{T} \right] = o_p(1).$$

For the third term, we can apply the WLLN for i.i.d sequence:

$$\frac{1}{T} \sum_{t=1}^T \epsilon_t^2 \xrightarrow{P} \sigma^2.$$

The final term, we see that

$$\frac{1}{T} \sum_{t=1}^T X_{t-1}\epsilon_t \xrightarrow{P} 0$$

whereby we applied the part 3 result proved earlier for the LLN for the AR(1) model.

By Slutsky theorem, we combine all these results to conclude that:

$$\frac{1}{T} \sum_{t=1}^T X_t^2 = \frac{\alpha^2}{1-\alpha^2} \left[\frac{X_0^2}{T} - \frac{X_T^2}{T} \right] + \frac{1}{T} \frac{\sum_{t=1}^T \epsilon_t^2}{1-\alpha^2} + \frac{1}{T} \frac{2\alpha}{1-\alpha^2} \sum_{t=1}^T X_{t-1} \epsilon_t \xrightarrow{P} 0 + \frac{\sigma^2}{1-\alpha^2} + 0$$

□

Theorem 483 (Central Limit Theorem for AR(1))

Let $X_t = \alpha X_{t-1} + \epsilon_t$, where $|\alpha| < 1$, the innovations are i.i.d with mean zero, variance σ^2 , and finite fourth moment, and X_0 satisfies the moment assumptions above. For any fixed non-negative integer j , the two central limit theorems for AR(1) processes are:

$$T^{-1/2} \sum_{t=1}^T X_{t-j} \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^2}{(1-\alpha)^2}\right). \quad (395)$$

$$T^{-1/2} \sum_{t=1}^T X_{t-1} \epsilon_t \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^4}{1-\alpha^2}\right). \quad (396)$$

Proof. We give the proof for the first result. First, recall the AR(1) model:

$$X_t = \alpha X_{t-1} + \epsilon_t$$

Then take sums:

$$\sum_{t=1}^T X_t = \alpha \sum_{t=1}^T X_{t-1} + \sum_{t=1}^T \epsilon_t.$$

Use our trick to rewrite the term as:

$$\sum_{t=1}^T X_t = \alpha \sum_{t=1}^T X_t + \alpha X_0 - \alpha X_T + \sum_{t=1}^T \epsilon_t.$$

Re-arrange:

$$(1-\alpha) \sum_{t=1}^T X_t = \alpha X_0 - \alpha X_T + \sum_{t=1}^T \epsilon_t.$$

and divide through:

$$\sum_{t=1}^T X_t = \frac{\alpha}{(1-\alpha)} X_0 - \frac{\alpha}{(1-\alpha)} X_T + \frac{1}{(1-\alpha)} \sum_{t=1}^T \epsilon_t.$$

Now scale by $T^{-1/2}$ to get:

$$\frac{1}{\sqrt{T}} \sum_{t=1}^T X_t = \frac{1}{\sqrt{T}} \frac{\alpha}{(1-\alpha)} X_0 - \frac{1}{\sqrt{T}} \frac{\alpha}{(1-\alpha)} X_T + \frac{1}{\sqrt{T}} \frac{1}{(1-\alpha)} \sum_{t=1}^T \epsilon_t.$$

We now analyse this term by term. For the first term on the right, we have that:

$$\frac{1}{\sqrt{T}} \frac{\alpha}{(1-\alpha)} X_0 = o_p(1).$$

For the second term, we have that:

$$\frac{1}{\sqrt{T}} \frac{\alpha}{(1-\alpha)} X_T = o(1) O_p(1) = o_p(1).$$

For the final term, we can use the regular central limit theorem for i.i.d sequence:

$$\frac{1}{\sqrt{T}} \frac{1}{(1-\alpha)} \sum_{t=1}^T \epsilon_t \xrightarrow{D} \frac{1}{1-\alpha} \mathcal{N}(0, \sigma^2) = \mathcal{N}\left(0, \frac{\sigma^2}{(1-\alpha)^2}\right).$$

Putting all this together, we conclude that:

$$T^{-1/2} \sum_{t=1}^T X_{t-j} \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^2}{(1-\alpha)^2}\right).$$

□

We now go back to our original AR(1) t-test. Let our AR(1) model with an intercept be:

$$Y_t = \alpha Y_{t-1} + \mu + \epsilon_t.$$

We can rewrite this as an AR(1) model without an intercept:

$$Y_t = \mu_Y + X_t$$

where

$$X_t = \alpha X_{t-1} + \epsilon_t$$

is an AR(1) model without an intercept. We now want to look at the test statistic to test the intercept:

$$H_0 : \mu_Y = 0.$$

The t-test statistic for this is:

$$t = \frac{\bar{Y} - 0}{\sqrt{\frac{s_Y^2}{T}}}$$

where the variance can be written as:

$$s_Y^2 = \frac{1}{T-1} \sum_{t=1}^T (Y_t - \bar{Y})^2$$

whereby

$$\bar{Y} = \frac{1}{T} \sum_{t=1}^T Y_t = \frac{1}{T} \sum_{t=1}^T (\mu_Y + X_t) = \mu_Y + \bar{X}.$$

Now, also note that:

$$\begin{aligned} T^{1/2}(\bar{Y} - \mu_Y) &= T^{1/2}\bar{X} \\ &= T^{-1/2} \sum_{t=1}^T X_t \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^2}{(1-\alpha)^2}\right) \end{aligned}$$

where we used the CLT for AR(1) model. We therefore have that the numerator converges to a normal distribution. Now, if we look the sample variance:

$$Y_t - \bar{Y} \leftrightarrow (\mu_Y + X_t) - (\mu_Y + \bar{X}) = X_t - \bar{X}.$$

Now, we look at the limiting result of s_Y^2 :

$$\begin{aligned} s_Y^2 &= \frac{1}{T-1} \sum_{t=1}^T \left(X_t - \bar{X} \right)^2 \\ &= \frac{T}{T-1} \left[\frac{1}{T} \sum_{t=1}^T X_t^2 - \left(\frac{1}{T} \sum_{t=1}^T X_t \right)^2 \right] \xrightarrow{P} \frac{\sigma^2}{1-\alpha^2}. \end{aligned}$$

where the first term inside the brackets uses the LLN for AR(1), and the second term uses the LLN for AR(1) plus the continuous mapping theorem.

Under $H_0 : \mu_Y = 0$, we therefore rewrite the t-statistic as:

$$\begin{aligned} t &= \frac{T^{1/2}\bar{Y}}{s_Y} \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^2/(1-\alpha)^2}{\sigma^2/(1-\alpha^2)}\right) \\ &= \mathcal{N}\left(0, \frac{1+\alpha}{1-\alpha}\right). \end{aligned}$$

where we again use Slutsky's theorem.

12.13 Higher Order Autoregression

12.13.1 AR(2) Model

So far, we have only studied the $AR(1)$ model:

$$Y_t = \alpha Y_{t-1} + \mu + \epsilon_t.$$

However, this is quite a restrictive assumption. We now aim to relax this assumption. We first look at the $AR(2)$ model.

Definition 484 (AR(2) With Intercept)

The $AR(2)$ with intercept model is given by:

$$Y_t = \alpha_1 Y_{t-1} + \alpha_2 Y_{t-2} + \mu + \epsilon_t \quad (397)$$

where μ is an intercept term, $\epsilon_t \sim \mathcal{N}(0, \sigma^2)$ and independent of $Y_{t-1}, \dots, Y_0$, and we also require two initial observations Y_{-1}, Y_0 .

Note that the $AR(2)$ model can also be written in vector notation:

$$Y_t = \boldsymbol{\alpha}^T \mathbf{Y}_{t-1} + \mu + \epsilon_t$$

where $\boldsymbol{\alpha} = (\alpha_1, \alpha_2)^T$ and $\mathbf{Y}_{t-1} = (Y_{t-1}, Y_{t-2})^T$. We can estimate such a model through the least squares estimator and in fact, coincides with the conditional MLE.

Now, notice that our $AR(2)$ model has an intercept term:

$$Y_t = \alpha_1 Y_{t-1} + \alpha_2 Y_{t-2} + \mu + \epsilon_t.$$

We want to rewrite this as an $AR(2)$ model without an intercept. The next proposition tells us how to do this.

Proposition 485 (AR(2) Model without an Intercept)

The $AR(2)$ model with an intercept can be written as an $AR(2)$ model without an intercept whereby Y_t is now the sum of a mean-zero process and its long-run mean μ_Y :

$$X_t = \alpha_1 X_{t-1} + \alpha_2 X_{t-2} + \epsilon_t \quad (398)$$

$$Y_t = X_t + \mu_Y \quad (399)$$

$$\mu_Y = \frac{\mu}{1 - \alpha_1 - \alpha_2} \quad (400)$$

Proof. First, subtract μ_Y from both sides:

$$Y_t - \mu_Y = \alpha_1 Y_{t-1} + \alpha_2 Y_{t-2} + \mu + \epsilon_t - \mu_Y$$

Now add and subtract $\alpha_1 \mu_Y$ and $\alpha_2 \mu_Y$ from the right hand side:

$$Y_t - \mu_Y = \alpha_1 Y_{t-1} + \alpha_2 Y_{t-2} + \mu + \epsilon_t - \mu_Y + \alpha_1 \mu_Y - \alpha_1 \mu_Y + \alpha_2 \mu_Y - \alpha_2 \mu_Y.$$

Now collect terms:

$$Y_t - \mu_Y = \alpha_1(Y_{t-1} - \mu_Y) + \alpha_2(Y_{t-2} - \mu_Y) + \left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y \right] + \epsilon_t$$

Now, note that $Y_t - \mu_Y$ will have mean zero. We define the mean zero process:

$$X_t \equiv Y_t - \mu_Y$$

to that our expression now becomes:

$$\begin{aligned} Y_t - \mu_Y &= \alpha_1(Y_{t-1} - \mu_Y) + \alpha_2(Y_{t-2} - \mu_Y) + \left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y \right] + \epsilon_t \\ \Leftrightarrow X_t &= \alpha_1 X_{t-1} + \alpha_2 X_{t-2} + \left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y \right] + \epsilon_t \end{aligned}$$

Hence, we have the expression now:

$$X_t = \alpha_1 X_{t-1} + \alpha_2 X_{t-2} + \left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y \right] + \epsilon_t$$

If we look carefully, nearly all the terms have mean zero:

$$\underbrace{X_t}_{\text{Mean Zero}} = \alpha_1 \underbrace{X_{t-1}}_{\text{Mean Zero}} + \alpha_2 \underbrace{X_{t-2}}_{\text{Mean Zero}} + \underbrace{\left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y \right]}_{(*)} + \underbrace{\epsilon_t}_{\text{Mean Zero}}.$$

We now look at the remaining $(*)$ term. Now, as the left-hand side has mean zero, and all the terms on the right-hand side has mean zero, since $(*)$ is a constant, we therefore require that it be equal to the zero in order for the expectation of $(*)$ to be zero. That is, we impose that:

$$\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y = 0.$$

We then re-arrange to solve for the long-run mean μ_Y to get:

$$\mu_Y = \frac{\mu}{1 - \alpha_1 - \alpha_2}.$$

Now, if we go back to:

$$X_t = \alpha_1 X_{t-1} + \alpha_2 X_{t-2} + \underbrace{\left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y \right]}_{(*)} + \epsilon_t$$

since we have imposed that $(*) = 0$, we get:

$$X_t = \alpha_1 X_{t-1} + \alpha_2 X_{t-2} + \epsilon_t$$

which is an AR(2) process without an intercept!

Hence, to summarise what we have done, we took an AR(2) model with an intercept:

$$Y_t = \alpha_1 Y_{t-1} + \alpha_2 Y_{t-2} + \mu + \epsilon_t$$

and rewrote it as an AR(2) model without an intercept whereby Y_t is now the sum of a mean-zero process and its

long-run mean:

$$X_t = \alpha_1 X_{t-1} + \alpha_2 X_{t-2} + \epsilon_t$$

$$Y_t = X_t + \mu_Y$$

$$\mu_Y = \frac{\mu}{1 - \alpha_1 - \alpha_2}$$

□

We have just shown how to write an AR(2) model with an intercept Y_t into an AR(2) model without an intercept X_t . Next, we want to be able to write this AR(2) model without an intercept as an AR(1) model without an intercept. Writing an AR(2) model in its **companion form** allows us to do this.

Definition 486 (Companion Form of AR(2))

Let X_t be an AR(2) model:

$$X_t = \alpha_1 X_{t-1} + \alpha_2 X_{t-2} + \epsilon_t$$

Then, the **companion form** of the AR(2) model can be written as:

$$\begin{bmatrix} X_t \\ X_{t-1} \end{bmatrix} = \begin{bmatrix} \alpha_1 & \alpha_2 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} X_{t-1} \\ X_{t-2} \end{bmatrix} + \begin{bmatrix} 1 \\ 0 \end{bmatrix} \epsilon_t \quad (401)$$

or in matrix form:

$$\mathbf{X}_t = \mathbf{A}\mathbf{X}_{t-1} + \tau\epsilon_t \quad (402)$$

whereby $\tau = (1, 0)^T$.

The companion form $\mathbf{X}_t$ now allows us to use similar techniques that we used for AR(1) processes. Recall that for an AR(1) model:

$$X_t = \alpha X_{t-1} + \epsilon_t$$

we could iteratively back-substitute to get the recursive form:

$$X_t = \sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0.$$

We can do something similar for the companion form $\mathbf{X}_t$ of the AR(2) model.

Proposition 487 (Recursive form of AR(2) Model)

Let $\mathbf{X}_t$ be the companion form of the AR(2) process X_t . Then, the recursive solution to $\mathbf{X}_t$ is:

$$\mathbf{X}_t = \sum_{j=0}^{t-1} \mathbf{A}^j \tau \epsilon_{t-j} + \mathbf{A}^t \mathbf{X}_0. \quad (403)$$

Jumping back to the AR(1) model:

$$X_t = \sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t X_0,$$

we showed that if the AR(1) possessed a stationarity solution ($|\alpha| < 1$) we can then take limits and under Kolmogorov 2 series criterion, write this as:

$$X_t = \sum_{j=0}^{\infty} \alpha^j \epsilon_{t-j}.$$

The question we are interested now in asking is that for our AR(2) companion form, can we write:

$$\mathbf{X}_t = \sum_{j=0}^{t-1} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j} + \mathbf{A}^t \mathbf{X}_0$$

in terms of an infinite moving average representation:

$$\mathbf{X}_t = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j}.$$

The answer is yes whereby we now need to check the condition that $\mathbf{A}$ is *smaller than 1*. We state what this means for matrices.

Proposition 488 (Existence of Stationary Solution for AR(2) Model)

Let $\mathbf{X}_t$ be the AR(2) in companion form with the companion matrix:

$$\mathbf{A} = \begin{bmatrix} \alpha_1 & \alpha_2 \\ 1 & 0 \end{bmatrix}.$$

Then, if all eigenvalues of $\mathbf{A}$ have modulus strictly less than 1 in absolute value, then the AR(2) model possesses a stationary solution.

Remark 489. *Technically speaking, this is the condition for stability which then implies stationarity under some additional assumptions.*

Remark 490. *We recall some facts from linear algebra. A matrix $\mathbf{A}$ is a linear function. We are interested in knowing which vectors are only stretched by the transformation $\mathbf{v}$. This stretching is determined by the eigenvalues λ . Hence, we want to know the vectors $\mathbf{v}$ such that:*

$$\mathbf{A}\mathbf{v} = \lambda\mathbf{v}$$

so that the vectors are only stretched. We can then rewrite this as:

$$(\mathbf{A} - \lambda\mathbf{I})\mathbf{v} = \mathbf{0}$$

This equation has a non-zero solution vector $\mathbf{v}$ if and only if

$$\det(\mathbf{A} - \lambda\mathbf{I}) = 0.$$

That is, the eigenvalues of $\mathbf{A}$ are the values λ that satisfy the equation:

$$\det(\mathbf{A} - \lambda\mathbf{I}) = 0.$$

By the fundamental theorem of algebra, we can write this as the product:

$$\det(\mathbf{A} - \lambda\mathbf{I}) = (\lambda_1 - \lambda)(\lambda_2 - \lambda)\dots(\lambda_n - \lambda) = 0$$

where n is the degree of the characteristic polynomial.

Recall that to compute the eigenvalues for the matrix $\mathbf{A}$:

$$\mathbf{A} = \begin{bmatrix} \alpha_1 & \alpha_2 \\ 1 & 0 \end{bmatrix}$$

you first compute the **characteristic polynomial** (recall that the roots to the characteristic polynomial are the eigenvalues) from computing the determinant of $\mathbf{A}$:

$$0 = \det(\mathbf{A} - \lambda \mathbf{I}) = (\alpha_1 - \lambda)(-\lambda) - \alpha_2 = \lambda^2 - \alpha_1\lambda - \alpha_2.$$

So we then check the roots for:

$$\lambda^2 - \alpha_1\lambda - \alpha_2 = \lambda^2 \left[1 - \alpha_1 \frac{1}{\lambda} - \alpha_2 \frac{1}{\lambda^2} \right] = \lambda^2 \alpha(1/\lambda) = 0$$

where $\alpha(z) = 1 - \alpha_1 z - \alpha_2 z^2$ is the autoregressive polynomial. As

$$\det(\mathbf{A} - \lambda \mathbf{I}) = \lambda^2 \alpha(1/\lambda) = 0$$

then this implies that the eigenvalues satisfy:

$$\alpha(1/\lambda) = 1 - \alpha_1 \frac{1}{\lambda} - \alpha_2 \frac{1}{\lambda^2} = 0.$$

If all roots λ that solves the characteristic polynomial has a modulus strictly less than one, then all eigenvalues of $\mathbf{A}$ are strictly less than one, and this means that the all entries of the matrix

$$\mathbf{A}^t \rightarrow 0$$

at an exponential rate.

Remark 491. Note that we can also equivalently express the condition for stationarity as:

$$1 - \alpha_1 z - \alpha_2 z^2 = 0$$

and check if all roots has modulus strictly **greater** than one.

Example 492

(Root Calculation). Suppose we had the autoregressive polynomial:

$$a(L) = (1 + 0.3L - 0.1L^2).$$

Now, we can factorise this:

$$(1 + 0.3L - 0.1L^2) = (1 + 0.5L)(1 - 0.2L) = 0.$$

Then, solving for roots, we have that:

$$1 + 0.5L = 0 \quad \leftrightarrow \quad L = -2$$

$$1 - 0.2L = 0 \quad \leftrightarrow \quad L = 5$$

Both roots have modulus greater than 1 so they are outside the unit circle and hence we have stability in the system.

Hence, going back to our AR(2) model:

$$\mathbf{X}_t = \sum_{j=0}^{t-1} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j} + \mathbf{A}^t \mathbf{X}_0$$

if all eigenvalues of $\mathbf{A}$ are strictly less than one, we will have that $\mathbf{A}^t \mathbf{X}_0 \rightarrow 0$ and in combination with Kolmogorov's 2 series criterion, we have that:

$$\mathbf{X}_t = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j}.$$

Proposition 493 (Moving Average Representation of AR(2))

Let $\mathbf{X}_t$ be the AR(2) model in companion form with the companion matrix $\mathbf{A}$. Suppose that all the eigenvalues of $\mathbf{A}$ has modulus strictly less than 1. Then, $\mathbf{X}_t$ can be represented in the infinite moving average form:

$$\mathbf{X}_t = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j} \quad (404)$$

where $\boldsymbol{\tau} = (1, 0)^T$.

Now that we have this moving average representation of the AR(2) model, we can now compute the statistical properties of this model just like we did for the AR(1) model.

Proposition 494 (Statistical Properties of AR(2) Model)

Let $\mathbf{X}_t$ be the AR(2) model in companion form. Suppose that $\mathbf{X}_t$ admits a stationary solution so that it has the infinite order representation:

$$\mathbf{X}_t = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j} \quad (405)$$

where $\boldsymbol{\tau} = (1, 0)^T$. Then, we have the following statistical properties of $\mathbf{X}_t$:

$$\mathbb{E}[\mathbf{X}_t] = 0 \quad (406)$$

$$\Sigma_X \equiv \text{Var}[\mathbf{X}_t] = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \sigma^2 \boldsymbol{\tau}^T (\mathbf{A}^j)^T \quad (407)$$

$$\text{Cov}[\mathbf{X}_t, \mathbf{X}_{t-s}] = \mathbf{A}^s \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \sigma^2 \boldsymbol{\tau}^T (\mathbf{A}^j)^T \quad (408)$$

We can in fact get a distribution for the AR(2) model.

Proposition 495 (Distribution of AR(2) Model)

Let $\mathbf{X}_t$ be the AR(2) model in companion form. Suppose that $\mathbf{X}_t$ admits a stationary solution so that it has the infinite order representation:

$$\mathbf{X}_t = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j} \quad (409)$$

where $\boldsymbol{\tau} = (1, 0)^T$ and $\epsilon_t \sim \mathcal{N}(0, \sigma^2)$ are i.i.d. Then, the distribution of X_t :

$$\mathbf{X}_t = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j} \sim \mathcal{N}(\mathbf{0}, \Sigma_X) \quad (410)$$

whereby $\Sigma_X \equiv \text{Var}[\mathbf{X}_t] = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \sigma^2 \boldsymbol{\tau}^T (\mathbf{A}^j)^T$.

12.13.2 AR(p) Model

We now generalise from the AR(2) model to the AR(p) model. Many of the results we derived for the AR(2) model hold for the AR(p) model.

Definition 496 (AR(p) Model with Intercept)

The AR(p) model with an intercept is given by:

$$Y_t = \sum_{j=1}^p \alpha_j Y_{t-j} + \mu + \epsilon_t \quad (411)$$

where μ is an intercept term, $\epsilon_t \sim \mathcal{N}(0, \sigma^2)$ and independent of $Y_{t-1}, \dots, Y_0$, and we also require p initial observations.

We show that we can write the AR(p) model with an intercept as an AR(p) model without an intercept just like we did for the AR(2) model.

Proposition 497 (AR(p) Model without an Intercept)

The AR(p) model with an intercept can be written as an AR(p) model without an intercept whereby Y_t is now the sum of a mean-zero process and its long-run mean μ_Y :

$$X_t = \alpha_1 X_{t-1} + \dots + \alpha_p X_{t-p} + \epsilon_t \quad (412)$$

$$Y_t = X_t + \mu_Y \quad (413)$$

$$\mu_Y = \frac{\mu}{1 - \alpha_1 - \dots - \alpha_p} \quad (414)$$

Proof. The proof is the same for AR(2) model except we subtract off more terms. First, subtract μ_Y from both sides:

$$Y_t - \mu_Y = \alpha_1 Y_{t-1} + \alpha_2 Y_{t-2} + \dots + \alpha_p Y_{t-p} + \mu + \epsilon_t - \mu_Y$$

Now add and subtract $\alpha_j \mu_Y$ for $j = 1, \dots, p$ from the right hand side:

$$Y_t - \mu_Y = \alpha_1 Y_{t-1} + \alpha_2 Y_{t-2} + \mu + \epsilon_t - \mu_Y + \alpha_1 \mu_Y - \alpha_1 \mu_Y + \alpha_2 \mu_Y - \alpha_2 \mu_Y + \dots + \alpha_p \mu_Y - \alpha_p \mu_Y$$

Now collect terms:

$$Y_t - \mu_Y = \alpha_1(Y_{t-1} - \mu_Y) + \alpha_2(Y_{t-2} - \mu_Y) + \dots + \alpha_p(Y_{t-p} - \mu_Y) + \left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y + \dots + \alpha_p\mu_Y \right] + \epsilon_t$$

Now, note that $Y_t - \mu_Y$ will have mean zero. We define the mean zero process:

$$X_t \equiv Y_t - \mu_Y$$

to that our expression now becomes:

$$\begin{aligned} Y_t - \mu_Y &= \alpha_1(Y_{t-1} - \mu_Y) + \alpha_2(Y_{t-2} - \mu_Y) + \dots + \alpha_p(Y_{t-p} - \mu_Y) + \left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y + \dots + \alpha_p\mu_Y \right] + \epsilon_t \\ \Leftrightarrow X_t &= \alpha_1X_{t-1} + \alpha_2X_{t-2} + \dots + \alpha_pX_{t-p} + \left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y + \dots + \alpha_p\mu_Y \right] + \epsilon_t \end{aligned}$$

Hence, we have the expression now:

$$X_t = \alpha_1X_{t-1} + \alpha_2X_{t-2} + \dots + \alpha_pX_{t-p} + \left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y + \dots + \alpha_p\mu_Y \right] + \epsilon_t$$

If we look carefully, nearly all the terms have mean zero:

$$\underbrace{X_t}_{\text{Mean Zero}} = \alpha_1 \underbrace{X_{t-1}}_{\text{Mean Zero}} + \alpha_2 \underbrace{X_{t-2}}_{\text{Mean Zero}} + \dots + \alpha_p \underbrace{X_{t-p}}_{\text{Mean Zero}} + \underbrace{\left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y + \dots + \alpha_p\mu_Y \right]}_{(*)} + \underbrace{\epsilon_t}_{\text{Mean Zero}}.$$

We now look at the remaining $(*)$ term. Now, as the left-hand side has mean zero, and all the terms on the right-hand side has mean zero, since $(*)$ is a constant, we therefore require that it be equal to the zero in order for the expectation of $(*)$ to be zero. That is, we impose that:

$$\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y + \dots + \alpha_p\mu_Y = 0.$$

We then re-arrange to solve for the long-run mean μ_Y to get:

$$\mu_Y = \frac{\mu}{1 - \alpha_1 - \alpha_2 - \dots - \alpha_p}.$$

Now, if we go back to:

$$X_t = \alpha_1X_{t-1} + \alpha_2X_{t-2} + \dots + \alpha_pX_{t-p} + \underbrace{\left[\mu - \mu_Y + \alpha_1\mu_Y + \alpha_2\mu_Y + \dots + \alpha_p\mu_Y \right]}_{(*)} + \epsilon_t$$

since we have imposed that $(*) = 0$, we get:

$$X_t = \alpha_1X_{t-1} + \alpha_2X_{t-2} + \dots + \alpha_pX_{t-p} + \epsilon_t$$

which is an AR(p) process without an intercept!

Hence, to summarise what we have done, we took an AR(p) model with an intercept:

$$Y_t = \alpha_1Y_{t-1} + \alpha_2Y_{t-2} + \dots + \alpha_pY_{t-p} + \mu + \epsilon_t$$

and rewrote it as an AR(p) model without an intercept whereby Y_t is now the sum of a mean-zero process and its

long-run mean:

$$X_t = \alpha_1 X_{t-1} + \alpha_2 X_{t-2} + \dots + \alpha_p X_{t-p} + \epsilon_t$$

$$Y_t = X_t + \mu_Y$$

$$\mu_Y = \frac{\mu}{1 - \alpha_1 - \alpha_2 - \dots - \alpha_p}$$

□

We have just shown how to write an AR(p) model with an intercept Y_t into an AR(p) model without an intercept X_t . Next, we want to be able to write this AR(p) model without an intercept as a vector AR(1) model without an intercept. Writing an AR(p) model in its **companion form** allows us to do this.

Definition 498 (Companion Form of AR(p))

Let X_t be an AR(p) model:

$$X_t = \alpha_1 X_{t-1} + \alpha_2 X_{t-2} + \dots + \alpha_p X_{t-p} + \epsilon_t$$

Then, the **companion form** of the AR(p) model can be written as:

$$\begin{bmatrix} X_t \\ X_{t-1} \\ \cdot \\ \cdot \\ \cdot \\ X_{t-p+2} \\ X_{t-p+1} \end{bmatrix} = \begin{bmatrix} \alpha_1 & \alpha_2 & \dots & \alpha_{p-1} & \alpha_p \\ 1 & 0 & \dots & 0 & 0 \\ 0 & 1 & \dots & 0 & 0 \\ 0 & 0 & \dots & 0 & 0 \\ 0 & 0 & \dots & 0 & 0 \\ 0 & 0 & \dots & 0 & 0 \\ 0 & 0 & \dots & 1 & 0 \end{bmatrix} \begin{bmatrix} X_{t-1} \\ X_{t-2} \\ \cdot \\ \cdot \\ \cdot \\ X_{t-p+1} \\ X_{t-p} \end{bmatrix} + \begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} \epsilon_t \quad (415)$$

or in matrix form:

$$\mathbf{X}_t = \mathbf{A}\mathbf{X}_{t-1} + \boldsymbol{\tau}\epsilon_t \quad (416)$$

whereby $\boldsymbol{\tau} = (1, 0, \dots, 0)^T$ is a $p \times 1$ vector.

We now want to write this matrix in a recursive form just like for the AR(1) model.

Proposition 499 (Recursive form of AR(p) Model)

Let $\mathbf{X}_t$ be the companion form of the AR(p) process X_t . Then, the recursive solution to $\mathbf{X}_t$ is:

$$\mathbf{X}_t = \sum_{j=0}^{t-1} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j} + \mathbf{A}^t \mathbf{X}_0. \quad (417)$$

We have written out the recursive form of the AR(p) model:

$$\mathbf{X}_t = \sum_{j=0}^{t-1} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j} + \mathbf{A}^t \mathbf{X}_0$$

. We want to now check if we can write it in terms of an infinite moving average representation:

$$\mathbf{X}_t = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j}.$$

To do so, we check the eigenvalues of $\mathbf{A}$.

Proposition 500 (Stability Condition for AR(p) Companion Matrix)

Let λ_j for $j = 1, \dots, p$ be the **eigenvalues** associated to the AR(p) companion matrix $\mathbf{A}$. Then, if:

$$\max_{j=1, \dots, p} |\lambda_j| < 1 \quad (418)$$

then the AR(p) model admits a stationarity solution.

Remark 501. Note that the eigenvalues of $\mathbf{A}$ solve the equation:

$$\det(\mathbf{A} - \lambda \mathbf{I}_p) = 0.$$

This is characterised by the characteristic polynomial:

$$\lambda^p - \alpha_1 \lambda^{p-1} - \dots - \alpha_p \lambda^0 = 0.$$

Then if all the roots to this characteristic equation lie **inside** the unit circle, then the stability condition holds.

Alternatively, we can set $\gamma = \frac{1}{\lambda}$ and write the characteristic polynomial as:

$$\gamma^0 - \alpha_1 \gamma^1 - \dots - \alpha_p \gamma^p = 0.$$

Then, if all the roots to this characteristic equation lie **outside** the unit circle, then the stability condition holds.

We now show how we could have derived this using lag polynomials. First, write the lag operator:

$$X_t = \alpha_1 X_{t-1} + \alpha_2 X_{t-2} + \dots + \alpha_p X_{t-p} + \epsilon_t$$

to get:

$$X_t = \alpha_1 L X_t + \alpha_2 L^2 X_t + \dots + \alpha_p L^p X_t + \epsilon_t$$

and re-arrange:

$$(1 - \alpha_1 L - \alpha_2 L^2 - \dots - \alpha_p L^p) X_t = \epsilon_t.$$

Define the autoregressive polynomial:

$$\alpha(L) = (1 - \alpha_1 L - \alpha_2 L^2 - \dots - \alpha_p L^p).$$

Now, we have that invertibility holds if the roots to:

$$(1 - \alpha_1 L - \alpha_2 L^2 - \dots - \alpha_p L^p)$$

are all greater than 1. This is equivalent to the eigenvalues of the companion matrix $\mathbf{A}$ being less than one. This allows us to invert the autoregressive polynomial to get the moving average representation.

We can now state the moving average representation of the AR(p) model.

Proposition 502 (Moving Average Representation of AR(p))

Let $\mathbf{X}_t$ be the AR(p) model in companion form with the companion matrix $\mathbf{A}$. Suppose that all the eigenvalues of $\mathbf{A}$ has modulus strictly less than 1. Then, $\mathbf{X}_t$ can be represented in the infinite moving average form:

$$\mathbf{X}_t = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j} \quad (419)$$

where $\boldsymbol{\tau} = (1, 0, \dots, 0)^T$ is a $p \times 1$ vector.

We can get the distribution for the AR(p) model as well.

Proposition 503 (Distribution of AR(p) Model)

Let $\mathbf{X}_t$ be the AR(p) model in companion form. Suppose that $\mathbf{X}_t$ admits a stationary solution so that it has the infinite order representation:

$$\mathbf{X}_t = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j} \quad (420)$$

where $\boldsymbol{\tau} = (1, 0, \dots, 0)^T$ is a $p \times 1$ vector and $\epsilon_t \sim \mathcal{N}(0, \sigma^2)$ are i.i.d. Then, the distribution of $\mathbf{X}_t$ is given by:

$$\mathbf{X}_t = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \epsilon_{t-j} \sim \mathcal{N}(\mathbf{0}, \Sigma_X) \quad (421)$$

whereby $\Sigma_X \equiv \text{Var}[\mathbf{X}_t] = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\tau} \sigma^2 \boldsymbol{\tau}^T (\mathbf{A}^j)^T$.

12.14 ADL Models

So far, we have looked at models whereby the variables of interested Y_t only depended on lag values of itself such as the AR(p) model:

$$Y_t = \sum_{j=1}^p \alpha_j Y_{t-j} + \epsilon_t.$$

We can now extend this case by now incorporating exogenous variables Z_t in the system. We now define such a model.

Definition 504 (Autoregressive Distributed Lag Model)

An **Autoregressive Distributed Lag** model (ADL(p,r)) model is given by:

$$Y_t = \sum_{j=1}^p \alpha_j Y_{t-j} + \sum_{j=0}^r \beta_j Z_{t-j} + \mu + \epsilon_t \quad (422)$$

where Z_t is an exogenous variable and μ is a deterministic term.

Remark 505. Notice that this is similar to an AR(p) model except we now have an exogenous variable Z_t as well. We get back an AR(p) model from an ADL(p,r) model when we set $\beta_j = 0$ for all $j = 0, \dots, r$.

We can estimate the parameters of the ADL(p,r) model via OLS. To interpret these ADL models, it is informative to write them in **equilibrium correct form**.

We now motivate equilibrium correction forms. First, consider the ADL(1,1) model:

$$Y_t = \omega Z_t + \alpha Y_{t-1} + \beta Z_{t-1} + \mu + \epsilon_t \quad (423)$$

whereby the series $\{Y_t\}_{t \in \mathbb{N}}$ and $\{Z_t\}_{t \in \mathbb{N}}$ are **stationary** with long-run means $\mathbb{E}[Y_t] = \mu_Y$ and $\mathbb{E}[Z_t] = \mu_Z$ and $\mathbb{E}[\epsilon_t] = 0$. If we shut down the noise process ϵ_t , then there is no noise in the system and in the **long-run** the system will converge to its long run equilibrium, which we denote by:

$$Y = \omega Z + \alpha Y + \beta Z + \mu$$

which we can then rewrite whereby we use the fact that Y_t and Z_t will converge to their long run mean due to stationarity of the series:

$$\mu_Y = (\omega + \beta)\mu_Z + \alpha\mu_Y + \mu. \quad (424)$$

Eq. (424) is the long equilibrium which our system Eq. (423) will converge to. Now, let us go back to Eq. (423):

$$Y_t = \omega Z_t + \alpha Y_{t-1} + \beta Z_{t-1} + \mu + \epsilon_t.$$

We now want to write this in equilibrium correction form. First, let us center Y_t and Z_t around their long-run means μ_Y and μ_Z by adding and subtracting μ_Y and $\omega\mu_Z$ from both sides:

$$\begin{aligned} Y_t - \mu_Y &= \omega Z_t + \alpha Y_{t-1} + \beta Z_{t-1} + \mu + \epsilon_t - \mu_Y + \omega\mu_Z - \omega\mu_Z \\ \Leftrightarrow (Y_t - \mu_Y) &= \omega(Z_t - \mu_Z) + \alpha Y_{t-1} + \beta Z_{t-1} + \mu + \epsilon_t - \mu_Y + \omega\mu_Z \end{aligned}$$

Now add and subtract $\alpha\mu_Y$ to the right hand side:

$$\begin{aligned} \Leftrightarrow (Y_t - \mu_Y) &= \omega(Z_t - \mu_Z) + \alpha Y_{t-1} + \beta Z_{t-1} + \mu + \epsilon_t - \mu_Y + \omega\mu_Z + \alpha\mu_Y - \alpha\mu_Y \\ \Leftrightarrow (Y_t - \mu_Y) &= \omega(Z_t - \mu_Z) + \alpha(Y_{t-1} - \mu_Y) + \beta Z_{t-1} + \mu + \epsilon_t - \mu_Y + \omega\mu_Z + \alpha\mu_Y \end{aligned}$$

Finally, add and subtract $\beta\mu_Z$ to the right hand side:

$$\begin{aligned} \Leftrightarrow (Y_t - \mu_Y) &= \omega(Z_t - \mu_Z) + \alpha(Y_{t-1} - \mu_Y) + \beta Z_{t-1} + \mu + \epsilon_t - \mu_Y + \omega\mu_Z + \alpha\mu_Y + \beta\mu_Z - \beta\mu_Z \\ \Leftrightarrow (Y_t - \mu_Y) &= \omega(Z_t - \mu_Z) + \alpha(Y_{t-1} - \mu_Y) + \beta(Z_{t-1} - \mu_Z) + \mu + \epsilon_t - \mu_Y + \omega\mu_Z + \alpha\mu_Y + \beta\mu_Z \end{aligned}$$

We now collect all the terms together and see we have:

$$(Y_t - \mu_Y) = \omega(Z_t - \mu_Z) + \alpha(Y_{t-1} - \mu_Y) + \beta(Z_{t-1} - \mu_Z) + \left[\mu - (1 - \alpha)\mu_Y + (\omega + \beta)\mu_Z \right] + \epsilon_t$$

Now, if you look carefully, we have that nearly all the terms will have mean zero:

$$\underbrace{(Y_t - \mu_Y)}_{\text{Mean 0}} = \omega \underbrace{(Z_t - \mu_Z)}_{\text{Mean 0}} + \alpha \underbrace{(Y_{t-1} - \mu_Y)}_{\text{Mean 0}} + \beta \underbrace{(Z_{t-1} - \mu_Z)}_{\text{Mean 0}} + \underbrace{\left[\mu - (1 - \alpha)\mu_Y + (\omega + \beta)\mu_Z \right]}_{(*)} + \underbrace{\epsilon_t}_{\text{Mean 0}}.$$

We now want to argue that the term (*) will also have mean zero. Notice that (*) is exactly the long-run equilibrium term Eq. (424) that we derived earlier but re-arranged:

$$\mu_Y = (\omega + \beta)\mu_Z + \alpha\mu_Y + \mu \quad \Leftrightarrow \quad 0 = (\omega + \beta)\mu_Z - (1 - \alpha)\mu_Y + \mu.$$

Hence, we can conclude that for the expression:

$$(Y_t - \mu_Y) = \omega(Z_t - \mu_Z) + \alpha(Y_{t-1} - \mu_Y) + \beta(Z_{t-1} - \mu_Z) + \underbrace{\left[\mu - (1 - \alpha)\mu_Y + (\omega + \beta)\mu_Z \right]}_{(*)} + \epsilon_t$$

that for the left hand side to equal zero, the term (*) must also equal zero as all other terms on the right hand side equal zero!

Now, go back to long-run relationship term (*):

$$\mu - (1 - \alpha)\mu_Y + (\omega + \beta)\mu_Z = 0$$

we can then solve for μ :

$$\mu = (1 - \alpha)\mu_Y - (\omega + \beta)\mu_Z \quad (425)$$

We can interpret Eq. (425) to be the long-run equilibrium relationship between the means μ_Y and μ_Z . That is, in the long-run, the means of Y_t and Z_t must satisfy this equilibrium relationship. We again rewrite Eq. (425) as:

$$0 = (1 - \alpha) \left(\mu_Y - \frac{\omega + \beta}{1 - \alpha} \mu_Z - \frac{\mu}{1 - \alpha} \right) \quad (426)$$

which is another way of saying what the long-run relationship between μ_Y and μ_Z is. We will soon see why Eq. (426) is useful.

We can now finally write the equilibrium correction form of:

$$Y_t = \omega Z_t + \alpha Y_{t-1} + \beta Z_{t-1} + \mu + \epsilon_t.$$

First, subtract Y_{t-1} from both sides and add and subtract ωZ_{t-1} on the right hand side to get:

$$\begin{aligned} Y_t - Y_{t-1} &= \omega Z_t + \alpha Y_{t-1} + \beta Z_{t-1} + \mu + \epsilon_t - Y_{t-1} + \omega Z_{t-1} - \omega Z_{t-1} \\ \Leftrightarrow Y_t - Y_{t-1} &= \omega(Z_t - Z_{t-1}) + (\alpha - 1)Y_{t-1} + (\beta + \omega)Z_{t-1} + \mu + \epsilon_t \\ \Leftrightarrow \Delta Y_t &= \omega \Delta Z_t - (1 - \alpha)Y_{t-1} + (\beta + \omega)Z_{t-1} + \mu + \epsilon_t \end{aligned}$$

where $\Delta = (1 - L)$ is the **difference operator** such that $\Delta Y_t = Y_t - Y_{t-1}$ and $\Delta Z_t = Z_t - Z_{t-1}$. Now, rewrite it as:

$$\begin{aligned} \Leftrightarrow \Delta Y_t &= \omega \Delta Z_t - (1 - \alpha)Y_{t-1} + (\beta + \omega)Z_{t-1} + \mu + \epsilon_t \\ \Leftrightarrow \Delta Y_t &= \omega \Delta Z_t - (1 - \alpha) \left[Y_{t-1} - \frac{(\beta + \omega)}{1 - \alpha} Z_{t-1} - \frac{\mu}{1 - \alpha} \right] + \epsilon_t \end{aligned}$$

We now have the new expression for our ADL(1,1) model as:

$$\Delta Y_t = \omega \Delta Z_t - (1 - \alpha) \underbrace{\left[Y_{t-1} - \frac{(\beta + \omega)}{1 - \alpha} Z_{t-1} - \frac{\mu}{1 - \alpha} \right]}_{(*)} + \epsilon_t.$$

But notice the term (*), it is actually Eq. (426) in the short-run! We compare the two expressions to make this

more clear:

$$\begin{bmatrix} Y_{t-1} - \frac{(\beta + \omega)}{1 - \alpha} Z_{t-1} - \frac{\mu}{1 - \alpha} \\ \mu_Y - \frac{\omega + \beta}{1 - \alpha} \mu_Z - \frac{\mu}{1 - \alpha} \end{bmatrix}$$

Therefore, we can conclude that in the long-run, (*) will be equal to 0:

$$(*) = (1 - \alpha) \left[Y_{t-1} - \frac{(\beta + \omega)}{1 - \alpha} Z_{t-1} - \frac{\mu}{1 - \alpha} \right] = 0.$$

This expression is called the **equilibrium correction mechanism**.

Definition 506 (Equilibrium Correction Mechanism)

The **equilibrium correction mechanism** for the ADL(1,1) model is:

$$ecm_{t-1} = \left[Y_{t-1} - \frac{(\beta + \omega)}{1 - \alpha} Z_{t-1} - \frac{\mu}{1 - \alpha} \right]. \quad (427)$$

We therefore replace this expression into our ADL(1,1) model:

$$\begin{aligned} \Delta Y_t &= \omega \Delta Z_t - (1 - \alpha) \left[Y_{t-1} - \frac{(\beta + \omega)}{1 - \alpha} Z_{t-1} - \frac{\mu}{1 - \alpha} \right] + \epsilon_t \\ \Delta Y_t &= \omega \Delta Z_t - (1 - \alpha) ecm_{t-1} + \epsilon_t \end{aligned}$$

If we look now at our ADL(1,1), we can see that the mean of each term is zero:

$$\underbrace{\Delta Y_t}_{\text{Mean zero}} = \omega \underbrace{\Delta Z_t}_{\text{Mean zero}} - (1 - \alpha) \underbrace{ecm_{t-1}}_{\text{Mean zero}} + \underbrace{\epsilon_t}_{\text{Mean zero}}$$

We can now officially state the equilibrium correction form.

Definition 507 (Equilibrium Correction Form)

Let the ADL(1,1) model be:

$$Y_t = \omega Z_t + \alpha Y_{t-1} + \beta Z_{t-1} + \mu + \epsilon_t.$$

Then, the equilibrium correction form for the ADL(1,1) model is given by:

$$\Delta Y_t = \omega \Delta Z_t - (1 - \alpha) ecm_{t-1} + \epsilon_t \quad (428)$$

$$ecm_{t-1} = \left[Y_{t-1} - \frac{(\beta + \omega)}{1 - \alpha} Z_{t-1} - \frac{\mu}{1 - \alpha} \right] \quad (429)$$

where $|\alpha| < 1$.

We now show the nice interpretation that comes with equilibrium correction form. The general intuition is that the model will adjust the process so that it tends to the long-run equilibrium. We look at two cases to illustrate this.

First, suppose that Y_{t-1} is large **relative** to Z_{t-1} . As a result, we see that the equilibrium correction mechanism term Eq. (429) will be positive: $ecm_{t-1} > 0$. This then feeds into the equilibrium correction form equation Eq. (428) via $-(1 - \alpha)ecm_{t-1}$ so that ΔY_t will be adjusted downwards. This means that Y_t is smaller as $\Delta Y_t = Y_t - Y_{t-1}$. Hence,

we can see that if Y_{t-1} is too big, then the equilibrium correction mechanism ecm_{t-1} will push next period Y_t down in order to bring it back in line with the long-run equilibrium.

Second, we do a similar exercise whereby we suppose Y_{t-1} is small **relative** to Z_{t-1} . Then, the equilibrium correction mechanism term Eq. (429) will be negative: $ecm_{t-1} < 0$. This then feeds into the equilibrium correction form equation Eq. (428) via $-(1-\alpha)ecm_{t-1}$ so that ΔY_t will be adjusted upwards. Therefore, if Y_{t-1} is too small, then the equilibrium correction mechanism ecm_{t-1} will push Y_t up in order to bring it back in line with the long-run equilibrium.

12.15 VAR Models

Recall that we have previously seen the AR(p) model:

$$y_t = \mu + \sum_{i=1}^p \phi_i y_{t-i} + \epsilon_t. \quad (430)$$

We now wish to generalise this through the **vector autoregressive model** (VAR model).

Definition 508 (Vector Autoregressive Model)

A VAR(p) of dimension k can be written as:

$$\mathbf{Y}_t = \boldsymbol{\mu} + \sum_{j=1}^p \Phi_j \mathbf{Y}_{t-j} + \boldsymbol{\epsilon}_t \quad (431)$$

where:

1. $\mathbf{Y}_t = (Y_{t,1}, \dots, Y_{t,k})^T$
2. $\boldsymbol{\mu} = (\mu_1, \dots, \mu_k)^T$
3. $\boldsymbol{\epsilon}_t = (\epsilon_{t,1}, \dots, \epsilon_{t,k})^T$
4. Φ_j is a $k \times k$ matrix of coefficients.

Furthermore, $\boldsymbol{\epsilon}_t \sim \mathcal{N}_k(\mathbf{0}, \Omega_k)$.

Remark 509. Notice that we have k variables and p lags.

Proposition 510 (Number of Conditional-Mean Parameters in VAR(p) Model)

The number of conditional-mean parameters in the VAR(p) model is $k + pk^2$.

Proof. We have k parameters in the intercept term $\boldsymbol{\mu}$. Furthermore, we have k^2 terms in the $k \times k$ matrix Φ_j . Then, we have p of these matrices so we need $p \times k^2$. $\square$

Remark 511. If the innovation covariance matrix is also counted, there are an additional $k(k+1)/2$ distinct parameters.

We have some interesting cases of the VAR(p) model.

Proposition 512 (Special Cases of VAR)

The special cases of the VAR(p) model are:

1. VAR(1) is when $p = 1$
2. AR(1) is when $p = 1$ and $k = 1$
3. AR(p) is when $k = 1$.

We now want to be able to write the VAR(p) model in a VAR(1) form by writing it in its companion form.

Proposition 513 (Companion Form of VAR(p) Model)

A VAR(p) model can be written in its **companion form** as a VAR(1) by writing it as:

$$\mathbf{X}_t = \mathbf{c} + \mathbf{A}\mathbf{X}_{t-1} + \mathbf{u}_t \quad (432)$$

whereby the **companion matrix** $\mathbf{A}$ is a $(kp \times kp)$ matrix and is given by:

$$\mathbf{A} = \begin{bmatrix} \Phi_1 & \Phi_2 & \dots & \Phi_{p-1} & \Phi_p \\ \mathbf{I}_k & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{I}_k & \dots & \mathbf{0} & \mathbf{0} \\ \cdot & \cdot & \dots & \cdot & \cdot \\ \mathbf{0} & \mathbf{0} & \dots & \mathbf{I}_k & \mathbf{0} \end{bmatrix} \quad (433)$$

and the vectors are:

$$\mathbf{X}_t = \left(\mathbf{Y}_t^T, \dots, \mathbf{Y}_{t-p+1}^T \right)^T \quad (kp \times 1) \quad (434)$$

$$\mathbf{u}_t = \left(\boldsymbol{\epsilon}_t^T, \mathbf{0}, \dots, \mathbf{0} \right)^T \quad (kp \times 1) \quad (435)$$

$$\mathbf{c} = \left(\boldsymbol{\mu}^T, \mathbf{0}, \dots, \mathbf{0} \right)^T \quad (kp \times 1) \quad (436)$$

$$(437)$$

Remark 514. Notice that in $\mathbf{X}_t$, each block $\mathbf{Y}_t^T$ is a $1 \times k$ row vector before the blocks are stacked and transposed.

We can state a similar condition for stationarity for the VAR(p) model as we did for the AR(p) model.

Proposition 515 (Stationarity for VAR(p) Model)

The VAR(p) possesses a stationary solution if all roots of the lag polynomial are **outside** the unit circle:

$$|\mathbf{I}_k - \sum_{j=1}^p \Phi_j z^j| = 0 \quad (438)$$

which only occurs for $|z| > 1$.

Equivalently, the VAR(p) possesses a stationary solution if all the eigenvalues of the companion matrix $\mathbf{A}$ lie **inside** the unit circle.

We can estimate a VAR(p) model's parameters by running ordinary least squares equation by equation on:

$$Y_t = \mu + \sum_{j=1}^p \Phi_j Y_{t-j} + \epsilon_t.$$

Alternatively, if we assume normality of the error terms, this is equivalent to maximum likelihood conditional on the initial values.

We can rewrite a VAR(1) with an intercept as a VAR(1) model without an intercept just like what we did for a AR(p) model. Note that notation wise, we have switched back to Y_t so don't think that this Y_t is the VAR(p) model without writing it in companion form, we just want to follow a similar notation to what we did for the AR(2) and AR(p) model.

Proposition 516 (VAR(1) Model Without an Intercept)

The k-dimensional VAR(1) model with an intercept:

$$Y_t = \mu + AY_{t-1} + \epsilon_t \quad (439)$$

can be written as a VAR(1) model without an intercept:

$$X_t = AX_{t-1} + \epsilon_t \quad (440)$$

$$Y_t = X_t + \mu_Y \quad (441)$$

$$\mu_Y = (I_k - A)^{-1}\mu \quad (442)$$

where μ_Y is the long-run mean.

Proof. The proof is the same as the AR(p) model. First, subtract the vector μ_Y from both sides:

$$Y_t - \mu_Y = \mu + AY_{t-1} + \epsilon_t - \mu_Y.$$

Now, add and subtract $A\mu_Y$ on the right-hand side:

$$Y_t - \mu_Y = \mu + AY_{t-1} + \epsilon_t - \mu_Y + A\mu_Y - A\mu_Y$$

and group terms to get:

$$\begin{aligned} Y_t - \mu_Y &= A(Y_{t-1} - \mu_Y) + \left[\mu + A\mu_Y - \mu_Y \right] + \epsilon_t \\ \Leftrightarrow Y_t - \mu_Y &= A(Y_{t-1} - \mu_Y) + \left[\mu - [I_k - A]\mu_Y \right] + \epsilon_t \end{aligned}$$

So we have:

$$Y_t - \mu_Y = A(Y_{t-1} - \mu_Y) + \left[\mu - [I_k - A]\mu_Y \right] + \epsilon_t$$

We then define our mean-zero process:

$$X_t \equiv Y_t - \mu_Y$$

so that we have:

$$X_t = AX_{t-1} + \underbrace{\left[\mu - [I_k - A]\mu_Y \right]}_{(*)} + \epsilon_t$$

Finally, we set (*) equal to zero:

$$\left[\boldsymbol{\mu} - [\mathbf{I}_k - \mathbf{A}]\boldsymbol{\mu}_Y \right] = 0 \quad \leftrightarrow \quad \boldsymbol{\mu}_Y = [\mathbf{I}_k - \mathbf{A}]^{-1}\boldsymbol{\mu}.$$

Hence, we end up with:

$$\mathbf{X}_t = \mathbf{A}\mathbf{X}_{t-1} + \boldsymbol{\epsilon}_t$$

which is a VAR(1) model without an intercept. □

We can then get the limiting distribution of the VAR(1) model.

Proposition 517 (Limiting Distribution of VAR(1) Model Without Intercept)

Let $\mathbf{X}_t$ be the k-dimensional VAR(1) model without an intercept:

$$\mathbf{X}_t = \mathbf{A}\mathbf{X}_{t-1} + \boldsymbol{\epsilon}_t \quad (443)$$

$$\mathbf{Y}_t = \mathbf{X}_t + \boldsymbol{\mu}_Y \quad (444)$$

$$\boldsymbol{\mu}_Y = (\mathbf{I}_k - \mathbf{A})^{-1}\boldsymbol{\mu} \quad (445)$$

$$\boldsymbol{\epsilon}_t \sim \mathcal{N}_k(\mathbf{0}, \Omega_k) \quad (446)$$

where $\boldsymbol{\mu}_Y$ is the long-run mean. Then, if $\mathbf{X}_t$ admits a stationary solution, then $\mathbf{X}_t$ has the limiting distribution:

$$\mathbf{X}_t | \mathbf{X}_0 \xrightarrow{D} \mathbf{X}_t^* = \sum_{j=0}^{\infty} \mathbf{A}^j \boldsymbol{\epsilon}_{t-j} \sim \mathcal{N}_k(\mathbf{0}, \Sigma_X) \quad (447)$$

whereby the moments are:

$$\begin{aligned} \mathbb{E}[\mathbf{X}_t^*] &= \mathbf{0} \\ \Sigma_X &= \sum_{j=0}^{\infty} \mathbf{A}^j \Omega_k (\mathbf{A}^j)^T. \end{aligned}$$

Correspondingly, the limiting distribution of $\mathbf{Y}_t$ is $\mathcal{N}_k(\boldsymbol{\mu}_Y, \Sigma_X)$.

Proposition 518 (Stationarity from Stability)

Suppose that y_t is a stable AR(p) process. Furthermore, suppose that the mean of the AR process does not change over time, the error term u_t has time-invariant variance σ_u^2 , and it's first and second moments are bounded. Then, the AR process y_t is stationary.

Remark 519. *Do not mix up stability and stationarity!*

Definition 520 (I(d) Process)

A stationary process is a $I(0)$ process. For $d \in \mathbb{N}$, a stochastic process y_t is called $I(d)$ if $\Delta^d y_t \equiv z_t$ is a stationary process with infinite-order moving average representation:

$$z_t = \sum_{j=0}^{\infty} \theta_j u_{t-j} = \theta(L) u_t \quad (448)$$

where the MA coefficients satisfy the conditions:

1. $\sum_{j=0}^{\infty} j|\theta_j| < \infty$
2. $\theta(1) = \sum_{j=0}^{\infty} \theta_j \neq 0$
3. u_t is a white noise process.

Definition 521 (Wold Representation Theorem)

Every K -dimensional non-deterministic zero mean stationary process y_t has an moving average representation:

$$y_t = \sum_{i=0}^{\infty} \Phi_i u_{t-i} \quad (449)$$

where $\Phi_0 = I_K$ and the coefficient matrices are square summable.

Let the K time series variable $y_t = (y_{1t}, \dots, y_{Kt})^T$ be decomposed as:

$$y_t = \mu_t + x_t \quad (450)$$

where μ_t is a deterministic part and x_t is a stochastic process with mean zero.

$$\mathbb{E}[y_t] = \mu_t.$$

Definition 522 (VAR(p) Model)

The stochastic process x_t is a linear VAR process of order p if it is of the form:

$$x_t = A_1 x_{t-1} + \dots + A_p x_{t-p} + u_t \quad (451)$$

where $x_t = (x_{1t}, \dots, x_{Kt})^T$ is a $(K \times 1)$ random vector, A_i are $K \times K$ coefficient matrices, and $u_t = (u_{1t}, \dots, u_{Kt})^T$ is a K -dimensional zero-mean white noise process with covariance matrix $\mathbb{E}[u_t u_t^T] = \Sigma_u$ such that $u_t \sim (\mathbf{0}, \Sigma_u)$.

Remark 523. We can relax the assumption of i.i.d errors u_t by requiring them to form a martingale difference sequence instead.

We can define the matrix polynomial in the lag operator:

$$A(L) = I_K - A_1 L - \dots - A_p L^p$$

This allows us to rewrite the VAR(p) model:

$$\begin{aligned} \mathbf{x}_t &= \mathbf{A}_1 \mathbf{x}_{t-1} + \dots + \mathbf{A}_p \mathbf{x}_{t-p} + \mathbf{u}_t \\ \Leftrightarrow \mathbf{x}_t - \mathbf{A}_1 \mathbf{x}_{t-1} - \dots - \mathbf{A}_p \mathbf{x}_{t-p} &= \mathbf{u}_t \\ \Leftrightarrow \left[\mathbf{I}_k - \mathbf{A}_1 L - \dots - \mathbf{A}_p L^p \right] \mathbf{x}_t &= \mathbf{u}_t \end{aligned}$$

Therefore, we have the expression:

$$\mathbf{A}(L) \mathbf{x}_t = \mathbf{u}_t \quad (452)$$

whereby $\mathbf{A}(L) = \mathbf{I}_k - \mathbf{A}_1 L - \dots - \mathbf{A}_p L^p$.

Recall that we were looking at the decomposition:

$$\mathbf{y}_t = \boldsymbol{\mu}_t + \mathbf{x}_t$$

Therefore, $\mathbf{y}_t$ inherits the stochastic properties of $\mathbf{x}_t$ after accounting for the deterministic component. In particular, stationarity of $\mathbf{y}_t$ requires $\boldsymbol{\mu}_t$ to be constant. To see the algebra more formally, let us multiply this expression by $\mathbf{A}(L)$:

$$\begin{aligned} \mathbf{y}_t &= \boldsymbol{\mu}_t + \mathbf{x}_t \\ \Leftrightarrow \mathbf{A}(L) \mathbf{y}_t &= \mathbf{A}(L) \boldsymbol{\mu}_t + \mathbf{A}(L) \mathbf{x}_t \\ \Leftrightarrow \mathbf{A}(L) \mathbf{y}_t &= \mathbf{A}(L) \boldsymbol{\mu}_t + \mathbf{u}_t \\ \Leftrightarrow \left[\mathbf{I}_k - \mathbf{A}_1 L - \dots - \mathbf{A}_p L^p \right] \mathbf{y}_t &= \mathbf{A}(L) \boldsymbol{\mu}_t + \mathbf{u}_t \\ \Leftrightarrow \mathbf{y}_t &= \mathbf{A}(L) \boldsymbol{\mu}_t + \mathbf{A}_1 \mathbf{y}_{t-1} + \dots + \mathbf{A}_p \mathbf{y}_{t-p} + \mathbf{u}_t \end{aligned}$$

This is called a **reduced form** as all the variables on the right hand side are lagged.

Definition 524 (Multivariate Stability Condition)

Let $\mathbf{y}_t = \boldsymbol{\mu}_t + \mathbf{x}_t$. Define the lag-operator associated to $\mathbf{x}_t$ as $\mathbf{A}(L) = \mathbf{I}_k - \mathbf{A}_1 L - \dots - \mathbf{A}_p L^p$. Then, the process $\mathbf{y}_t$ is stable if the stochastic process $\mathbf{x}_t$ is stable, which holds when all the roots of the characteristic polynomial are outside the complex unit circle:

$$\det(\mathbf{A}(z)) = \det(\mathbf{I}_k - \mathbf{A}_1 z - \dots - \mathbf{A}_p z^p) \neq 0 \quad (453)$$

for all $|z| \leq 1$.

Proposition 525 (Stationarity of Multivariate Process)

Let $\mathbf{y}_t = \boldsymbol{\mu} + \mathbf{x}_t$, where $\boldsymbol{\mu}$ is constant. If $\mathbf{x}_t$ is stable and has white-noise innovations with finite variance, then $\mathbf{y}_t$ is covariance stationary.

Definition 526 (Jordan Canonical Form)

Let $\mathbf{A}$ be a $m \times m$ matrix with eigenvalues $\lambda_1, \dots, \lambda_n$. Then, there exists a non-singular matrix $\mathbf{P}$ such that:

$$\mathbf{A} = \mathbf{P} \begin{bmatrix} \mathbf{\Lambda}_1 & \dots & 0 \\ \dots & \dots & \dots \\ 0 & \dots & \mathbf{\Lambda}_n \end{bmatrix} \mathbf{P}^{-1} \equiv \mathbf{P} \mathbf{\Lambda} \mathbf{P}^{-1} \quad (454)$$

where:

$$\mathbf{\Lambda}_i = \begin{bmatrix} \lambda_i & 1 & 0 & \dots & 0 \\ 0 & \lambda_i & 1 & \dots & 0 \\ \dots & \dots & \dots & \dots & \dots \\ 0 & 0 & 0 & \dots & 1 \\ 0 & 0 & 0 & \dots & \lambda_i \end{bmatrix} \quad (455)$$

For any matrix $\mathbf{A}$, we can decompose it into its Jordan canonical form. The following result is why we care about it.

Proposition 527 (Convergence of a Matrix with Eigenvalues Inside the Unit Circle)

Suppose that $\mathbf{A}$ is a $m \times m$ matrix with eigenvalues $\lambda_1, \dots, \lambda_n$ which all have modulus less than 1, $|\lambda_i| < 1$ for $i = 1, \dots, n$. Then

$$\mathbf{A}^j \rightarrow 0 \quad (456)$$

as $j \rightarrow \infty$.

Proof. First, we can write $\mathbf{A}$ into its Jordan canonical form:

$$\mathbf{A} = \mathbf{P} \mathbf{\Lambda} \mathbf{P}^{-1}.$$

This then implies that:

$$\mathbf{A}^j = (\mathbf{P} \mathbf{\Lambda} \mathbf{P}^{-1})^j = \mathbf{P} \mathbf{\Lambda}^j \mathbf{P}^{-1}.$$

Write each Jordan block as $\mathbf{\Lambda}_i = \lambda_i \mathbf{I} + \mathbf{N}_i$, where $\mathbf{N}_i$ is nilpotent of order r_i . Then

$$\mathbf{\Lambda}_i^j = \sum_{q=0}^{r_i-1} \binom{j}{q} \lambda_i^{j-q} \mathbf{N}_i^q.$$

For every fixed q , the polynomial term $\binom{j}{q}$ is dominated by the geometric decay of $|\lambda_i|^{j-q}$. Thus every Jordan block converges to zero, so $\mathbf{\Lambda}^j \rightarrow \mathbf{0}$ and hence $\mathbf{A}^j \rightarrow \mathbf{0}$.

□

12.15.1 Granger Causality

A variable y_{2t} is called Granger causal for y_{1t} if the information in the past and present values of y_{2t} helps reduce the expectation of the squared prediction error for y_{1t} .

Definition 528 (Granger Causality)

Let Ω_t be the information set containing all information relevant to predicting y_{1t} . Denote the optimal h -step prediction of y_{1t} at origin t based on the information set in Ω_t by $y_{1,t+h|\Omega}$ and the corresponding mean-squared prediction error by $\sigma_{y_1}^2(h|\Omega_t)$. Then the process y_{2t} is said to **Granger cause** y_{1t} if:

$$\sigma_{y_1}^2(h|\Omega_t) < \sigma_{y_1}^2(h|\Omega_t \setminus \{y_{2s}|s \leq t\}) \quad (457)$$

for at least one $h \in \{1, 2, \dots\}$ where $\Omega_t \setminus \{y_{2s}|s \leq t\}$ denotes the set of all relevant information in the universe apart from the past and present information about the y_{2t} process.

y_{2t} Granger causes y_{1t} if y_{1t} can be predicted with a lower mean squared error by taking into the account of the information in y_{2s} for $s \leq t$.

We describe how to test Granger causality in a VAR(p) process.

$$\begin{bmatrix} y_{1t} \\ y_{2t} \end{bmatrix} = \begin{bmatrix} \nu_1 \\ \nu_2 \end{bmatrix} + \sum_{i=1}^p \begin{bmatrix} a_{11,i} & a_{12,i} \\ a_{21,i} & a_{22,i} \end{bmatrix} \begin{bmatrix} y_{1,t-i} \\ y_{2,t-i} \end{bmatrix} + \begin{bmatrix} u_{1t} \\ u_{2t} \end{bmatrix}. \quad (458)$$

Then, y_{2t} does not Granger cause y_{1t} if and only if $a_{12,i} = 0$ for $i = 1, 2, \dots, p$. We can test this using Wald tests.

There are a few issues with Granger causality. First, y_{2t} may Granger cause y_{1t} if y_{3t} is present but not Granger cause if y_{3t} is not present. Another issue of Granger causality is that it does not look at the instantaneous relationships of the variables.

We can also decide how many lags to include in our VAR model.

Definition 529 (Information Criterion)

The information criterion has the form:

$$C(m) = \log(\det(\tilde{\Sigma}_u(m))) + c_T \phi(m) \quad (459)$$

where m is the candidate lag order, $\phi(m)$ is a function of order m that penalizes large larger orders, c_T is a sequence of weights, and $\tilde{\Sigma}_u(m) = T^{-1} \sum_{t=1}^T \hat{u}_t \hat{u}_t^T$ is the residual covariance matrix estimator for a reduced-form VAR model of order m .

We are choosing the number of lags m that the VAR(m) model should have. Since we have mK lagged regressor in each equation and K (variables) equations in the VAR model, the total number of regressors (and hence parameters) is: $\phi(m) = mK^2$ If we include an intercept, this then becomes: $\phi(m) = mK^2 + K$.

AIC is not a consistent estimator of the lag length, whereas BIC and HQC (Hannan–Quinn) are consistent.

12.16 Misspecification Testing in Time Series

12.17 Unit Root Testing

So far, we have looked at models where the coefficient $|\alpha| < 1$. We now are interested in the case to see what happens when $\alpha = 1$. We shall soon see that under such a situation, we require new asymptotic results.

Definition 530 (Unit Root)

We say that the process:

$$Y_t = \alpha Y_{t-1} + \epsilon_t \quad (460)$$

possesses a **unit root** if $\alpha = 1$.

Recall that previously, through back-substitution, we could write:

$$Y_t = \alpha Y_{t-1} + \epsilon_t$$

as the sum of error terms:

$$Y_t = \sum_{j=0}^{t-1} \alpha^j \epsilon_{t-j} + \alpha^t Y_0$$

where Y_0 is the initial value of the series. Now, if we have that $\alpha = 1$, this then becomes:

$$Y_t = \sum_{j=0}^{t-1} \epsilon_{t-j} + Y_0. \quad (461)$$

The reason why this is significant is stated in the next proposition.

Proposition 531 (Unit Roots Processes are Not Weakly Stationary)

An AR(1) series that possesses a unit root is **not weakly stationary**.

Proof. Write the AR(1) model as:

$$Y_t = \sum_{j=0}^{t-1} \epsilon_{t-j} + Y_0.$$

Now compute the variance:

$$Var[Y_t] = Var\left[\sum_{j=0}^{t-1} \epsilon_{t-j}\right] + Var[Y_0] = \sigma^2 t + Var[Y_0]$$

We can see that the variance now depends on the time index t . Therefore, it is not weakly stationary. $\square$

This result is in fact more general whereby any process with a unit root is **not weakly stationary**. An issue is that when we construct OLS estimates, we usually require the time series process to be stationary, but now with a nonstationary process, this may produce invalid OLS estimates. We can also show that we have issues with doing inference and hypothesis testing.

Suppose we had the AR(1) model:

$$Y_t = \alpha Y_{t-1} + \epsilon_t$$

where ϵ_t is i.i.d with $\mathbb{E}[\epsilon_t] = 0$ and $Var[\epsilon_t] = \sigma^2$. We then recall that by the law of large numbers for AR(1):

$$T^{-1} \sum_{t=1}^T Y_{t-1}^2 \xrightarrow{P} \frac{\sigma^2}{1 - \alpha^2}$$

and by the central limit theorem for AR(1):

$$T^{-1/2} \sum_{t=1}^T Y_{t-1} \epsilon_t \xrightarrow{D} \mathcal{N}\left(0, \frac{\sigma^4}{1-\alpha^2}\right).$$

Now, if we were interested in testing $\hat{\alpha}$, we can construct the test statistic:

$$\sqrt{T}(\hat{\alpha} - \alpha) = \frac{T^{-1/2} \sum_{t=1}^T Y_{t-1} \epsilon_t}{T^{-1} \sum_{t=1}^T Y_{t-1}^2} \xrightarrow{D} \frac{1}{\frac{\sigma^2}{1-\alpha^2}} \mathcal{N}\left(0, \frac{\sigma^4}{1-\alpha^2}\right) = \mathcal{N}(0, 1-\alpha^2).$$

This derivation assumes $|\alpha| < 1$, so we cannot obtain the unit-root limit by substituting $\alpha = 1$ into the stationary formula. At the boundary, the numerator and denominator have different stochastic orders: $\hat{\alpha} - 1$ converges at rate T , rather than $\sqrt{T}$, and has a non-Gaussian limiting distribution. We derive that distribution below.

We now will need to introduce some new terms in order to show what our new limiting distribution will be in the presence of unit roots.

Definition 532 (Random Walk)

A random walk is defined as:

$$Y_t = \sum_{s=1}^t \epsilon_s. \quad (462)$$

where ϵ_t is i.i.d with $\mathbb{E}[\epsilon_t] = 0$ and $\mathbb{E}[\epsilon_t^2] = \sigma^2$.

Remark 533. Another name for $\sum_{s=1}^t \epsilon_s$ is a **stochastic trend**.

To see why this is called a random walk, notice that if we difference this:

$$Y_t - Y_{t-1} = \sum_{s=1}^t \epsilon_s - \sum_{s=1}^{t-1} \epsilon_s = \epsilon_t$$

which is an independent random variable of previous terms. Hence, where the name random walk comes from.

We also introduce another tool that will be quite useful for us.

Definition 534 (Brownian Motion)

A Brownian motion $(B_u)_{u \in [0,1]}$ is characterized by:

1. $B_0 = 0$
2. it has independent increments
3. $B_t - B_s \sim \mathcal{N}(0, t-s)$ for $0 \leq s < t \leq 1$
4. $u \rightarrow B_u$ is continuous (continuous sample paths).

With this, we can state some new asymptotic results that will be quite useful.

Proposition 535 (AR(1) with Unit Root Asymptotic Results)

Suppose that $Y_t = Y_{t-1} + \epsilon_t$ with i.i.d ϵ_t with $\mathbb{E}[\epsilon_t] = 0$ and $\mathbb{E}[\epsilon_t^2] = \sigma^2$. Then,

$$\frac{1}{T^2} \sum_{t=1}^T Y_{t-1}^2 \xrightarrow{D} \sigma^2 \int_0^1 B_u^2 du \quad (463)$$

$$\frac{1}{T} \sum_{t=1}^T Y_{t-1} \epsilon_t \xrightarrow{D} \sigma^2 \int_0^1 B_u dB_u = \frac{1}{2} (B_1^2 - 1) \sigma^2 \quad (464)$$

$$\frac{1}{T} \sum_{t=1}^T \epsilon_t^2 \xrightarrow{P} \sigma^2 \quad (465)$$

where $(B_u)_{u \in [0,1]}$ is a Brownian motion.

We will use these results to help us find the limiting distribution of test statistics under the null hypothesis of a unit root. In particular, we want to find the limiting distribution of:

$$T(\hat{\alpha} - 1). \quad (466)$$

Theorem 536 (Limiting Distribution of Test Statistic Under Unit Root)

The limiting distribution of the test statistic under an unit root is:

$$T(\hat{\alpha} - 1) \xrightarrow{D} \frac{\int_0^1 B_u dB_u}{\int_0^1 B_u^2 du}. \quad (467)$$

Proof. We have that:

$$\begin{aligned} T(\hat{\alpha} - 1) &= T\hat{\pi} \\ &\stackrel{(1)}{=} \frac{\frac{1}{T} \sum_{t=1}^T Y_{t-1} \epsilon_t}{\frac{1}{T^2} \sum_{t=1}^T Y_{t-1}^2} \xrightarrow{D} \frac{\sigma^2 \int_0^1 B_u dB_u}{\sigma^2 \int_0^1 B_u^2 du} = \frac{\int_0^1 B_u dB_u}{\int_0^1 B_u^2 du}. \end{aligned}$$

where $\stackrel{(1)}{=}$ uses the formula of the least squares estimator. Furthermore, the numerator uses Eq. (464) and the denominator uses Eq. (463). $\square$

Remark 537. Notice that the scaling factor is now T rather than $\sqrt{T}$. As a result, we say that $\hat{\alpha}$ is **super-consistent** in the presence of a unit root.

We now discuss how we test for a *unit root* $H_0 : \alpha = 1$

$$Y_t = \alpha Y_{t-1} + \epsilon_t.$$

Notice that is equivalent to testing: $H_0 : \pi = 0$ in:

$$\Delta Y_t = \pi Y_{t-1} + \epsilon_t$$

as $\pi = \alpha - 1$.

We now state an important distribution that will come up in our testing.

Definition 538 (Dickey-Fuller Distribution)

The Dickey Fuller distribution is given by:

$$DF = \frac{\int_0^1 B_u dB_u}{(\int_0^1 B_u^2 du)^{1/2}} \quad (468)$$

where B_u is Brownian motion.

We can now state the distribution of our test.

Proposition 539 (Distribution of Test Statistic for Unit Root Testing)

Let Y_t be an AR(1) process. Suppose we are testing $H_0 : \pi = 0$ in:

$$\Delta Y_t = \pi Y_{t-1} + \epsilon_t$$

Then, the limiting distribution of the test statistic is:

$$Z_{\pi=0} = \frac{\hat{\pi}}{\hat{\sigma}_{\hat{\pi}}} \xrightarrow{D} DF \quad (469)$$

where DF is the Dickey-Fuller distribution.

Proof. First, write out:

$$\begin{aligned} Z_{\pi=0} &= \frac{\hat{\pi}}{\hat{\sigma}_{\hat{\pi}}} \\ &= \frac{\sum_{t=1}^T Y_{t-1} \epsilon_t}{\sum_{t=1}^T Y_{t-1}^2 \cdot \hat{\sigma} (\sum_{t=1}^T Y_{t-1}^2)^{-1/2}} \\ &= \frac{\sum_{t=1}^T Y_{t-1} \epsilon_t}{(\sum_{t=1}^T Y_{t-1}^2)^{1/2} \cdot \hat{\sigma}} \\ &= \frac{\frac{1}{T} \sum_{t=1}^T Y_{t-1} \epsilon_t}{\frac{1}{T} (\sum_{t=1}^T Y_{t-1}^2)^{1/2} \cdot \hat{\sigma}} \xrightarrow{D} \frac{\sigma^2 \int_0^1 B_u dB_u}{\sigma^2 (\int_0^1 B_u^2 du)^{1/2}} = \frac{\int_0^1 B_u dB_u}{(\int_0^1 B_u^2 du)^{1/2}} \end{aligned}$$

that is $\hat{\sigma} \rightarrow \sigma$. □

We are testing:

$$H_0 : \pi = 0 \quad \text{vs} \quad H_1 : \pi \in (-2, 0)$$

whereby we reject if $Z_{\pi=0}$ falls below the 5% quantile of the DF distribution. That is, we reject for small values of the test statistic:

$$Z_{\pi=0} = \frac{\hat{\pi}}{\hat{\sigma}_{\hat{\pi}}} = \frac{(\hat{\pi} - \pi)}{\hat{\sigma}_{\hat{\pi}}} \quad \text{under } H_0 : \pi = 0.$$

This is a one-sided test whereby we only reject for small values of the test statistic. That is, it is a left hand side test.

12.17.1 Deterministic Terms

We have derived the distribution under the null hypothesis of a unit root in the absence of deterministic terms. We now want to see the effects including a constant term and a trend term.

First suppose that we suspect that an intercept is relevant (usually in the case when we think the mean of Y_t is non-zero). That is, the model is:

$$Y_t = \alpha Y_{t-1} + \mu_c + \epsilon_t$$

with the unit-root null hypothesis:

$$H_0 : \alpha = 1,$$

where μ_c is included as an unrestricted nuisance parameter.

Under H_0 , we end up with a random walk:

$$Y_t = Y_{t-1} + \mu_c + \epsilon_t.$$

If $\mu_c \neq 0$, this is a random walk with drift and its recursive formulation contains a deterministic linear trend:

$$Y_t = Y_0 + \sum_{s=1}^t \epsilon_s + \mu_c t$$

where $\mu_c t$ is a deterministic trend.

We can represent this model with an intercept in a different manner:

$$\begin{aligned} Y_t &= \alpha Y_{t-1} + \mu_c + \epsilon_t \\ \Leftrightarrow Y_t - Y_{t-1} &= \alpha Y_{t-1} - Y_{t-1} + \mu_c + \epsilon_t \\ \Leftrightarrow \Delta Y_t &= (\alpha - 1)Y_{t-1} + \mu_c + \epsilon_t \end{aligned}$$

Then, the unit root hypothesis becomes:

$$H_0 : \alpha = 1 \quad \Leftrightarrow \quad H_0 : \pi = 0,$$

with μ_c unrestricted. Under this null, we get:

$$\Delta Y_t = \mu_c + \epsilon_t$$

and therefore solving for the level process gives:

$$Y_t = Y_0 + \mu_c t + \sum_{s=1}^t \epsilon_s.$$

The limiting distribution of the t-statistic is different from the no-intercept case:

$$Z_{\pi=0} \xrightarrow{D} DF_C$$

where DF_C is a different distribution to the ordinary Dickey Fuller distribution we saw before.

We could separately test the joint hypothesis $H_0 : \pi = \mu_c = 0$ using a likelihood-ratio test statistic:

$$LR_{\pi=\mu_c=0} \xrightarrow{D} DF_C^{(2)}$$

where $DF_C^{(2)}$ is another different distribution.

We now look at a different case where we now include a trend in our model:

$$\Delta Y_t = \pi Y_{t-1} + \mu_c + \mu_I t + \epsilon_t$$

with the null hypothesis:

$$H_0 : \pi = 0,$$

where μ_c and μ_I are included as unrestricted nuisance parameters for the Dickey–Fuller t-test.

The limiting distribution of the t-statistic based on the least squares estimator is:

$$Z_{\pi=0} \xrightarrow{D} DF_{C+I}$$

where DF_{C+I} is the Dickey–Fuller distribution appropriate when a constant and linear trend are included. We could also test the joint hypothesis $H_0 : \pi = \mu_I = 0$, with μ_c unrestricted, using a likelihood-ratio test statistic:

$$LR_{\pi=\mu_I=0} \xrightarrow{D} DF_{C+I}^{(2)}$$

where again $DF_{C+I}^{(2)}$ is a different distribution.

12.18 Cointegration

12.18.1 Beveridge-Nelson Decomposition

Suppose that we have an AR(p) process:

$$y_t = \alpha_1 y_{t-1} + \dots + \alpha_p y_{t-p} + \epsilon_t.$$

We can rewrite this using lag operators:

$$\begin{aligned} y_t &= \alpha_1 y_{t-1} + \dots + \alpha_p y_{t-p} + \epsilon_t \\ \Leftrightarrow y_t - \alpha_1 y_{t-1} - \dots - \alpha_p y_{t-p} &= \epsilon_t \\ \Leftrightarrow y_t(1 - \alpha_1 L - \dots - \alpha_p L^p) &= \epsilon_t \end{aligned}$$

Define the autoregressive polynomial:

$$\alpha(L) = (1 - \alpha_1 L - \alpha_2 L^2 - \dots - \alpha_p L^p).$$

Now, we have that invertibility holds if the roots to:

$$(1 - \alpha_1 L - \alpha_2 L^2 - \dots - \alpha_p L^p)$$

are all greater than 1. We can rewrite this as:

$$(1 - \alpha_1 L - \alpha_2 L^2 - \dots - \alpha_p L^p) = (1 - \lambda_1 L) \dots (1 - \lambda_p L)$$

where λ_i are the reciprocals of the roots to the polynomial. A **unit root** is when a root $\frac{1}{\lambda_i}$ equals one.

Definition 540 (Integrated of Order d)

A univariate process with d unit roots is called **integrated of order d I(d)**.

If a process has a single unit root, then the process is I(1). Then, we can construct a stable process by taking first differences of the original process:

$$\Delta y_t = y_t - y_{t-1}.$$

If a process is I(d), it can be made stable by differencing d times:

$$\Delta^d y_t = (1 - L)^d y_t.$$

Definition 541 (I(1) Process)

A process y_t is an I(1) process if

$$\Delta y_t = w_t$$

where w_t is a stationary process with infinite MA representation:

$$w_t = \sum_{j=0}^{\infty} \theta_j u_{t-j}.$$

A unit root process is an autoregression with a single root at unity. This means that the differenced process ΔY_t

is serially correlated but stationary. The Beveridge Nelson decomposition is a way to decompose a unit process into a permanent (random walk) component and a transitory (stationary) component.

If ΔY_t is stationary, then by the **Wold representation** theorem, it has the form:

$$\Delta Y_t = b(L)\epsilon_t \quad (470)$$

where

$$b(L) = \sum_{j=0}^{\infty} b_j L^j.$$

The lag polynomial can be generalised as:

$$b(z) = \sum_{j=0}^{\infty} b_j z^j.$$

Note that:

$$b(1) = \sum_{j=0}^{\infty} b_j.$$

To proceed, we need to show a useful property of lag polynomials.

Proposition 542 (Factorisation of Lag Polynomial)

The lag polynomial $b(z)$ can be factorised as:

$$b(z) = b(1) + (1 - z)b^*(z) \quad (471)$$

where $b(1) = \sum_{j=0}^{\infty} b_j$ and $b^*(z)$ is the lag polynomial:

$$b^*(z) = \sum_{j=0}^{\infty} b_j^* z^j$$

$$b_j^* = - \sum_{l=j+1}^{\infty} b_l$$

Proof. To show this, we work from the right hand side and show that it equals the left hand side:

$$b(z) = b(1) + (1 - z)b^*(z).$$

Therefore, we have that:

$$\begin{aligned}
b(1) + (1 - z)b^*(z) &= \sum_{j=0}^{\infty} b_j - \sum_{j=0}^{\infty} \sum_{l=j+1}^{\infty} b_l z^j (1 - z) \\
&= \sum_{j=0}^{\infty} b_j - \sum_{j=0}^{\infty} \sum_{l=j+1}^{\infty} b_l z^j + \sum_{j=0}^{\infty} \sum_{l=j+1}^{\infty} b_l z^{j+1} \\
&= \sum_{j=0}^{\infty} b_j - \sum_{j=1}^{\infty} b_j - \sum_{j=1}^{\infty} \sum_{l=j+1}^{\infty} b_l z^j + \sum_{j=0}^{\infty} \sum_{l=j+1}^{\infty} b_l z^{j+1} \\
&= b_0 - \sum_{j=1}^{\infty} \sum_{l=j+1}^{\infty} b_l z^j + \sum_{j=1}^{\infty} \sum_{l=j}^{\infty} b_l z^j \\
&= b_0 + \sum_{j=1}^{\infty} b_j z^j \\
&= \sum_{j=0}^{\infty} b_j z^j \\
&= b(z).
\end{aligned}$$

□

Now going back to our original problem, we had the stationary ΔY_t in a Wold representation:

$$\Delta Y_t = b(L)\epsilon_t.$$

We can apply our factorisation to $b(L)$ and this gives us the Beveridge–Nelson decomposition.

Theorem 543 (Beveridge–Nelson Decomposition)

Let Y_t be a unit root process. Then, the **Beveridge–Nelson decomposition** is given by:

$$Y_t = S_t + U_t + V_0 \quad (472)$$

where $S_t = \sum_{i=1}^t \xi_i$ is a random walk, whereby ξ_i is a white noise process, U_t is a strictly stationary process, and V_0 is an initial condition.

Proof.

$$\Delta Y_t = b(L)\epsilon_t.$$

We then apply our factorisation to $b(L) = b(1) + (1 - L)b^*(L)$ to get:

$$\begin{aligned}
\Delta Y_t &= \left(b(1) + (1 - L)b^*(L) \right) \epsilon_t \\
&= b(1)\epsilon_t + (1 - L)b^*(L)\epsilon_t \\
&= b(1)\epsilon_t + b^*(L)\epsilon_t - b^*(L)\epsilon_{t-1} \\
&\equiv \xi_t + U_t - U_{t-1}
\end{aligned}$$

where ξ_t is the permanent innovation:

$$\xi_t = b(1)\epsilon_t$$

and U_t, U_{t-1} is strictly stationary under assumptions of **square-summability** of b_j^* . We then sum both sides and use telescoping sums:

$$\begin{aligned}\sum_{i=1}^t (Y_i - Y_{i-1}) &= \sum_{i=1}^t \xi_i + \sum_{i=1}^t (U_i - U_{i-1}) \\ \Leftrightarrow Y_t - Y_0 &= \sum_{i=1}^t \xi_i + U_t - U_0 \\ \Leftrightarrow Y_t &= \sum_{i=1}^t \xi_i + U_t + Y_0 - U_0 \\ &\equiv Y_t = S_t + U_t + V_0\end{aligned}$$

where we defined:

$$\begin{aligned}S_t &= \sum_{i=1}^t \xi_i \\ U_t &= b^*(L)\epsilon_t \\ V_0 &= Y_0 - U_0\end{aligned}$$

□

If we look at:

$$Y_t = S_t + U_t + V_0$$

we have that S_t is the **permanent** (trend) component of Y_t as it determines the long-run behaviour of Y_t and U_t is the **transitory component**.

12.18.2 Cointegration

We may suspect that there are equilibrium relationships between economic variables. Let us collect k variables into:

$$\mathbf{y}_t = \begin{bmatrix} y_{1t} \\ \dots \\ \dots \\ \dots \\ y_{Kt} \end{bmatrix}.$$

Suppose that the K variables have that their long-run equilibrium relation is:

$$\boldsymbol{\beta}^T \mathbf{y}_t = \beta_1 y_{1t} + \dots + \beta_K y_{Kt} = 0$$

where $\boldsymbol{\beta} = (\beta_1, \dots, \beta_K)^T$. However, in any given period, they may instead be at the deviation from the long-run level:

$$\boldsymbol{\beta}^T \mathbf{y}_t = z_t$$

where z_t measures the deviation from equilibrium. However, it may be the case that all the variables are driven by a **common stochastic trend** so that all the variables in $\mathbf{y}_t$ are $I(1)$. However, it may be the case that there exists a

linear combination:

$$z_t \equiv \beta^T y_t$$

such that z_t is $I(0)$. Then we say that $y_t \sim CI(1, 1)$ whereby the vector β is called a **cointegrating/cointegration vector** and y_t is called a **cointegrated process**.

Definition 544 (Cointegrated of Order 1)

A set of $I(1)$ time series variables y_t is **cointegrated** if there exists a linear combination $\beta^T y_t$ of these variables such that:

$$\beta^T y_t \quad (473)$$

is $I(0)$. β is called the cointegrated vector.

Remark 545. Another different definition of cointegration is that a K -dimensional process y_t is integrated of order d , $y_t \sim I(d)$ if $\Delta^d y_t$ is stable and $\Delta^{d-1} y_t$ is not stable.

Example 546

Suppose that $y_t \sim I(1)$ and $z_t \sim I(1)$ and

$$2y_t - 3z_t \sim I(0).$$

Then (y_t, z_t) is cointegrated with cointegration vector $\beta = (2, -3)$.

We now give a motivation of error correction models and see their relationship to cointegration. In error correction models, changes in a variable depends on the deviations from some equilibrium relationship. We now show an example of how does cointegration and ECM relate to each other.

Example 547

(Basic ECM Cointegration). Suppose that y_{1t} and y_{2t} is the price of the same good in two different markets. Suppose that the equilibrium relationship between the two variables are $y_{1t} = \beta_1 y_{2t}$. By the error correction model, the change in the level of y_{1t} and y_{2t} depends on the deviations from this equilibrium in the period $t - 1$:

$$\Delta y_{1t} = \alpha_1 (y_{1,t-1} - \beta_1 y_{2,t-1}) + u_{1t}$$

$$\Delta y_{2t} = \alpha_2 (y_{1,t-1} - \beta_1 y_{2,t-1}) + u_{2t}$$

To further extend upon this, suppose that Δy_{it} also depends on previous changes in both variables too:

$$\Delta y_{1t} = \alpha_1 (y_{1,t-1} - \beta_1 y_{2,t-1}) + \gamma_{11,1} \Delta y_{1,t-1} + \gamma_{12,1} \Delta y_{2,t-1} + u_{1t}$$

$$\Delta y_{2t} = \alpha_2 (y_{1,t-1} - \beta_1 y_{2,t-1}) + \gamma_{21,1} \Delta y_{1,t-1} + \gamma_{22,1} \Delta y_{2,t-1} + u_{2t}$$

Now, suppose that y_{1t} and y_{2t} are $I(1)$ variables and u_{it} are white noise errors (therefore they are stable). We have that every term except for the first term on the right hand side is a stable process. We re-arrange for this term:

$$\alpha_1 (y_{1,t-1} - \beta_1 y_{2,t-1}) = \Delta y_{1t} - \gamma_{11,1} \Delta y_{1,t-1} - \gamma_{12,1} \Delta y_{2,t-1} - u_{1t}$$

$$\alpha_2 (y_{1,t-1} - \beta_1 y_{2,t-1}) = \Delta y_{2t} - \gamma_{21,1} \Delta y_{1,t-1} - \gamma_{22,1} \Delta y_{2,t-1} - u_{2t}$$

For the right-hand side to be stable, it must be the case that the left-hand side $y_{1,t-1} - \beta_1 y_{2,t-1}$ is also stable, assuming that $\alpha_1 \neq 0$ and $\alpha_2 \neq 0$. However, by definition, this is cointegration! Therefore, we see that the error correction mechanism is also the source of the cointegration.

12.18.3 Cointegration for VAR(2)

We now generalise our example to a VAR(2) model.

Example 548

(ECMT Cointegration with VAR(2)). We take the same price example from before:

$$\begin{aligned}\Delta y_{1t} &= \alpha_1(y_{1,t-1} - \beta_1 y_{2,t-1}) + \gamma_{11,1} \Delta y_{1,t-1} + \gamma_{12,1} \Delta y_{2,t-1} + u_{1t} \\ \Delta y_{2t} &= \alpha_2(y_{1,t-1} - \beta_1 y_{2,t-1}) + \gamma_{21,1} \Delta y_{1,t-1} + \gamma_{22,1} \Delta y_{2,t-1} + u_{2t}\end{aligned}$$

Now, we rewrite this in terms of vectors and matrices:

$$\mathbf{y}_t - \mathbf{y}_{t-1} = \boldsymbol{\alpha} \boldsymbol{\beta}^T \mathbf{y}_{t-1} + \boldsymbol{\Gamma}_1 (\mathbf{y}_{t-1} - \mathbf{y}_{t-2}) + \mathbf{u}_t$$

whereby

$$\boldsymbol{\alpha} \equiv \begin{bmatrix} \alpha_1 \\ \alpha_2 \end{bmatrix} \quad \boldsymbol{\beta}^T \equiv \begin{bmatrix} 1 & -\beta_1 \end{bmatrix} \quad \boldsymbol{\Gamma}_1 \equiv \begin{bmatrix} \gamma_{11,1} & \gamma_{12,1} \\ \gamma_{21,1} & \gamma_{22,1} \end{bmatrix}.$$

We then re-arrange to get:

$$\mathbf{y}_t = \boldsymbol{\alpha} \boldsymbol{\beta}^T \mathbf{y}_{t-1} + \mathbf{y}_{t-1} + \boldsymbol{\Gamma}_1 \mathbf{y}_{t-1} - \boldsymbol{\Gamma}_1 \mathbf{y}_{t-2} + \mathbf{u}_t$$

Now, we notice that we now have a VAR(2) representation:

$$\mathbf{y}_t = (\mathbf{I}_K + \boldsymbol{\Gamma}_1 + \boldsymbol{\alpha} \boldsymbol{\beta}^T) \mathbf{y}_{t-1} - \boldsymbol{\Gamma}_1 \mathbf{y}_{t-2} + \mathbf{u}_t$$

Hence, as we have shown that our cointegrated system can be written as a VAR(2) model.

We have shown how to go from cointegration to VAR(2). We now want to do the reverse.

Let us generalise the coefficient matrices in order to tie this to unit roots. Suppose we had the VAR(2) process:

$$\mathbf{y}_t = \mathbf{A}_1 \mathbf{y}_{t-1} + \mathbf{A}_2 \mathbf{y}_{t-2} + \mathbf{u}_t$$

where $\mathbf{y}_t$ is a K-dimensional vector now. Rewrite this using lag operator:

$$(\mathbf{I}_K - \mathbf{A}_1 L - \mathbf{A}_2 L^2) \mathbf{y}_t = \mathbf{u}_t$$

Suppose that the process is unstable whereby the autoregressive polynomial:

$$\det(\mathbf{I}_K - \mathbf{A}_1 L - \mathbf{A}_2 L^2) = (1 - \lambda_1 L) \dots (1 - \lambda_n L)$$

has at least one root being equal to one. Elaborating on this, recall that to find a given root j , we use:

$$1 - \lambda_j L = 0 \quad \leftrightarrow \quad L = \frac{1}{\lambda_j}.$$

Meanwhile, all other roots are outside the unit circle. Since we have a unit root, this implies that:

$$\det(\mathbf{I}_K - \mathbf{A}_1 L - \mathbf{A}_2 L^2) = (1 - \lambda_1 L) \dots (1 - \lambda_n L) = 0$$

and therefore the matrix:

$$\boldsymbol{\Pi} \equiv \mathbf{A}_1 + \mathbf{A}_2 - \mathbf{I}_K$$

is singular as it has a determinant of 0. Suppose that the $\text{rank}(\boldsymbol{\Pi}) = r < K$. Then we can decompose this matrix:

$$\boldsymbol{\Pi} = \boldsymbol{\alpha} \boldsymbol{\beta}^T$$

where α and β are $(K \times r)$ matrices. As we assumed $I(1)$ variables, we can difference the VAR(2):

$$\begin{aligned}
y_t &= A_1 y_{t-1} + A_2 y_{t-2} + u_t \\
\leftrightarrow y_t - y_{t-1} &= A_1 y_{t-1} + A_2 y_{t-2} + u_t - y_{t-1} \\
\leftrightarrow y_t - y_{t-1} &= A_1 y_{t-1} + A_2 y_{t-2} + u_t - y_{t-1} + A_2 y_{t-1} - A_2 y_{t-1} \\
\leftrightarrow y_t - y_{t-1} &= (A_1 + A_2 - I_K) y_{t-1} - A_2 (y_{t-1} - y_{t-2}) + u_t \\
\leftrightarrow y_t - y_{t-1} &= (A_1 + A_2 - I_K) y_{t-1} - A_2 \Delta y_{t-1} + u_t \\
\leftrightarrow \Delta y_t &= \Pi y_{t-1} + \Gamma_1 \Delta y_{t-1} + u_t
\end{aligned}$$

where we defined $\Gamma_1 \equiv -A_2$ and $\Pi \equiv A_1 + A_2 - I_K$. We re-arrange to get:

$$\Delta y_t - \Gamma_1 \Delta y_{t-1} - u_t = \Pi y_{t-1}.$$

As all the terms on the left-hand side are stationary, this implies that $\Pi y_{t-1} = \alpha \beta^T y_{t-1}$ is also stationary. Furthermore, stationarity holds when you multiply by $(\alpha^T \alpha)^{-1} \alpha^T$:

$$(\alpha^T \alpha)^{-1} \alpha^T \alpha \beta^T y_{t-1} = \beta^T y_{t-1}$$

so then this implies that $\beta^T y_{t-1}$ is also stationary. But this is the definition of a cointegrating relationship again!

12.18.4 Cointegration for VAR(p)

We now look at the K-dimensional case. Suppose that we have the K-dimensional VAR(p) process:

$$y_t = A_1 y_{t-1} + \dots + A_p y_{t-p} + u_t. \quad (474)$$

where each variable is either $I(1)$ or $I(0)$.

Definition 549 (Cointegration of VAR(p))

Let y_t be a K-dimensional VAR(p) process:

$$y_t = A_1 y_{t-1} + \dots + A_p y_{t-p} + u_t. \quad (475)$$

Then, y_t is called **cointegrated of rank r** if

$$\Pi \equiv -(I_K - A_1 - \dots - A_p) \quad (476)$$

has rank r. Then, the matrix Π can be written as:

$$\Pi = \alpha \beta^T \quad (477)$$

where α and β are $(K \times r)$ matrices with rank r, so β^T is $(r \times K)$. The columns of β are called the **cointegrating vectors** and α is called the **loading matrix**.

Remark 550. Notice the two special cases where $r = 0$ and $r = K$. When $r = 0$, then Δy_t has a stable VAR(p-1) representation. When $r = K$, then $\det(-\Pi) \neq 0$ and hence the VAR(p) has no unit roots so that y_t is a stable VAR(p) process.

Given our VAR(p) model:

$$\mathbf{y}_t = \mathbf{A}_1 \mathbf{y}_{t-1} + \dots + \mathbf{A}_p \mathbf{y}_{t-p} + \mathbf{u}_t$$

we can write this VAR into a vector error correction form.

Definition 551 (Vector Error Correction Model)

Let $\mathbf{y}_t$ be a VAR(p) process. Then, the associated **vector error correction model** for $\mathbf{y}_t$ is:

$$\Delta \mathbf{y}_t = \mathbf{\Pi} \mathbf{y}_{t-1} + \mathbf{\Gamma}_1 \Delta \mathbf{y}_{t-1} + \dots + \mathbf{\Gamma}_{p-1} \Delta \mathbf{y}_{t-p+1} + \mathbf{u}_t \quad (478)$$

where

$$\mathbf{\Pi} = -(\mathbf{I}_k - \mathbf{A}_1 - \dots - \mathbf{A}_p)$$

$$\mathbf{\Gamma}_i = -(\mathbf{A}_{i+1} + \dots + \mathbf{A}_p)$$

for $i = 1, \dots, p-1$.

Remark 552. Notice that due to the decomposition of $\mathbf{\Pi}$, we could have written the VECM as:

$$\Delta \mathbf{y}_t = \boldsymbol{\alpha} \boldsymbol{\beta}^T \mathbf{y}_{t-1} + \mathbf{\Gamma}_1 \Delta \mathbf{y}_{t-1} + \dots + \mathbf{\Gamma}_{p-1} \Delta \mathbf{y}_{t-p+1} + \mathbf{u}_t$$

We have shown how to go from a VAR(p) to a VECM. We can go in reverse too. Suppose we are given a VECM:

$$\Delta \mathbf{y}_t = \mathbf{\Pi} \mathbf{y}_{t-1} + \mathbf{\Gamma}_1 \Delta \mathbf{y}_{t-1} + \dots + \mathbf{\Gamma}_{p-1} \Delta \mathbf{y}_{t-p+1} + \mathbf{u}_t$$

$$\mathbf{\Pi} = -(\mathbf{I}_k - \mathbf{A}_1 - \dots - \mathbf{A}_p)$$

$$\mathbf{\Gamma}_i = -(\mathbf{A}_{i+1} + \dots + \mathbf{A}_p)$$

we can also retrieve the VAR(p) form by setting:

$$\mathbf{A}_1 = \mathbf{\Pi} + \mathbf{I}_k + \mathbf{\Gamma}_1$$

$$\mathbf{A}_i = \mathbf{\Gamma}_i - \mathbf{\Gamma}_{i-1} \quad \text{for } i = 2, \dots, p-1$$

$$\mathbf{A}_p = -\mathbf{\Gamma}_{p-1}$$

which gives us back the VAR(p) model:

$$\mathbf{y}_t = \mathbf{A}_1 \mathbf{y}_{t-1} + \dots + \mathbf{A}_p \mathbf{y}_{t-p} + \mathbf{u}_t$$

Linear Algebra Review

We collect some results from linear algebra.

Definition 553 (Orthogonal Complement)

If U is a subset of vector space V , then the **orthogonal complement** of U , denoted $U_\perp$, is the set of all vectors in V that are orthogonal to every vector in U :

$$U_\perp = \{v \in V : v \cdot u = 0 \text{ for every } u \in U\}. \quad (479)$$

Proposition 554 (Direct Sum)

Suppose U is a finite-dimensional subspace of the vector space V . Then

$$V = U \oplus U_{\perp} \quad (480)$$

where $U_{\perp}$ is the orthogonal complement to U .

This result is important as this means that for any vector $v \in V$, we can decompose it into the sum of a vector in U and in a vector in $U_{\perp}$.

Proposition 555 (Dimension of the Orthogonal Complement)

Suppose V is finite-dimensional and U is a subspace of V . Then

$$\dim U_{\perp} = \dim V - \dim U. \quad (481)$$

Example 556

Let the vector space be $V = \mathbb{R}^5$. Then, suppose we have the matrix:

$$\mathbf{A} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \\ 2 & 6 \\ 3 & 9 \\ 5 & 3 \end{bmatrix}.$$

We can then define the subset $W \subset V$ as:

$$W = \{u \in V : \mathbf{A}x = u, x \in \mathbb{R}^2\}.$$

Notice that $\mathbf{A}$ has a column rank of 2. Now, we can define the orthogonal complement as:

$$W_{\perp} = \{v \in V : u \cdot v = 0, u \in W\}.$$

Alternatively, we can construct an orthogonal matrix $\tilde{\mathbf{A}}$:

$$\tilde{\mathbf{A}} = \begin{bmatrix} -2 & -3 & -5 \\ -6 & -9 & -3 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}.$$

Here, we have that:

$$W_{\perp} = \{v \in V : \tilde{\mathbf{A}}y = v, y \in \mathbb{R}^3\}.$$

Notice that every column in $\mathbf{A}$ is orthogonal to every column in $\tilde{\mathbf{A}}$.

We now summarise this section of linear algebra. Let $\mathbf{X}$ be a $m \times n$ matrix whereby $m \geq n$. Then, we denote $\mathbf{X}_{\perp}$ as the **orthogonal complement** of the $m \times n$ matrix with $\text{rank}(\mathbf{X}) = n$. Then, $\mathbf{X}_{\perp}$ is a $(m \times (m - n))$ matrix with

$\text{rank}(\mathbf{X}_\perp) = m - n$. Furthermore, this matrix is orthogonal:

$$\mathbf{X}^T \mathbf{X}_\perp = 0.$$

12.18.5 Granger Representation Theorem

We now give another representation of a cointegrated system $\mathbf{y}_t$. To do so, we need two assumptions.

Assumption 557 (Granger Representation)

Let $\mathbf{y}_t$ be a cointegrated vector autoregressive model.

1. The characteristic polynomial of $\mathbf{y}_t$ has stationary and unit roots only.
2. The number of unit roots equal $K - r$.

Definition 558 (Cointegrating Rank)

Let $\mathbf{\Pi}$ be the matrix associated to the VECM of $\mathbf{y}_t$. Then $\mathbf{y}_t$ has cointegrating rank r if:

$$\mathbf{\Pi} = \boldsymbol{\alpha}\boldsymbol{\beta}^T \quad (482)$$

where $\boldsymbol{\alpha}, \boldsymbol{\beta} \in \mathbb{R}^{K \times r}$.

Theorem 559 (Granger VECM Representation Theorem)

Let $\mathbf{y}_t$ be a K -dimensional cointegrated $I(1)$ process with cointegrating rank r where $0 \leq r < K$. Furthermore, assume that $\mathbf{y}_t$ has the vector error correction model representation:

$$\Delta \mathbf{y}_t = \mathbf{\Pi} \mathbf{y}_{t-1} + \boldsymbol{\Gamma}_1 \Delta \mathbf{y}_{t-1} + \dots + \boldsymbol{\Gamma}_{p-1} \Delta \mathbf{y}_{t-p+1} + \mathbf{u}_t \quad (483)$$

Then, the **Granger representation** is given by:

$$\mathbf{y}_t = \mathbf{C} \sum_{i=1}^t \mathbf{u}_i + \mathbf{C}^*(L) \mathbf{u}_t + \mathbf{y}_0^* \quad (484)$$

where

$$\mathbf{C} = \boldsymbol{\beta}_\perp \left[\boldsymbol{\alpha}_\perp^T \left(\mathbf{I}_K - \sum_{i=1}^{p-1} \boldsymbol{\Gamma}_i \right) \boldsymbol{\beta}_\perp \right]^{-1} \boldsymbol{\alpha}_\perp^T$$

$\mathbf{C} \in \mathbb{R}^{K \times K}$, has rank $K - r$, and

$$\mathbf{C}^*(L) \mathbf{u}_t = \sum_{j=0}^{\infty} \mathbf{C}_j^* \mathbf{u}_{t-j}$$

is an $I(0)$ process with $\mathbf{C}_j^*$ is a coefficient matrix and $\mathbf{y}_0^*$ contains the initial values. Furthermore, it follows that $\mathbf{y}_t$ is a cointegrated $I(1)$ process with cointegrating vectors $\boldsymbol{\beta}$. Finally, we have that:

$$\boldsymbol{\beta}^T \mathbf{C} = 0 \quad \mathbf{C} \boldsymbol{\alpha} = 0$$

where $(\boldsymbol{\alpha}, \boldsymbol{\beta})$ has dimensions $K \times r$ and the orthogonal complements $(\boldsymbol{\alpha}_\perp, \boldsymbol{\beta}_\perp)$ has dimensions $K \times (K - r)$.

Proof. (Sketch). The matrix:

$$\alpha_{\perp}^T \left(I_K - \sum_{i=1}^{p-1} \Gamma_i \right) \beta_{\perp}$$

is a $(K - r) \times (K - r)$ matrix and hence for this matrix to be invertible, it must be full rank $K - r$. This then implies that $\mathbf{C}$ has rank $K - r$. The number of unit roots is exactly $K - r$ as a result of this too. $\square$

Remark 560. *The Granger representation theorem is a multivariate version of the Beveridge-Nelson decomposition. Here, $\mathbf{y}_t$ is decomposed into $I(1)$ and $I(0)$ components. In particular, the process $\mathbf{y}_t$ is driven by $K - r$ $I(1)$ components and r $I(0)$ components.*

The first term:

$$\mathbf{C} \sum_{i=1}^t \mathbf{u}_i$$

consists of K random walks $\sum_{i=1}^t \mathbf{u}_i$ being multiplied by a matrix with rank $K - r$, denoted by $\mathbf{C}$. This means that there are in fact $K - r$ stochastic trends driving the system.

The second term:

$$\mathbf{C}^*(L) \mathbf{u}_t = \sum_{j=0}^{\infty} \mathbf{C}_j^* \mathbf{u}_{t-j}$$

is an $I(0)$ component where $\mathbf{C}_j^$ is a coefficient matrix similar to that seen in the Beveridge-Nelson decomposition. The expression for this coefficient matrix is complicated and not needed.*

This representation has some interesting properties. First, if we multiply the matrix by the cointegrating vector β , we eliminate the random walk and we get:

$$\begin{aligned} \mathbf{y}_t &= \mathbf{C} \sum_{i=1}^t \mathbf{u}_i + \mathbf{C}^*(L) \mathbf{u}_t + \mathbf{y}_0^* \\ \beta^T \mathbf{y}_t &= \beta^T \mathbf{C} \sum_{i=1}^t \mathbf{u}_i + \beta^T \mathbf{C}^*(L) \mathbf{u}_t + \beta^T \mathbf{y}_0^* \\ \beta^T \mathbf{y}_t &= \beta^T \left[\beta_{\perp} \left[\alpha_{\perp}^T \left(I_K - \sum_{i=1}^{p-1} \Gamma_i \right) \beta_{\perp} \right]^{-1} \alpha_{\perp}^T \right] \sum_{i=1}^t \mathbf{u}_i + \beta^T \mathbf{C}^*(L) \mathbf{u}_t + \beta^T \mathbf{y}_0^* \\ \beta^T \mathbf{y}_t &= \beta^T \mathbf{C}^*(L) \mathbf{u}_t + \beta^T \mathbf{y}_0^* \end{aligned}$$

and since $\mathbf{C}^*(L) \mathbf{u}_t$ is an $I(0)$ process, we are left with an $I(0)$ process once we premultiply $\mathbf{y}_t$ by β^T .

13 Panel Data

We now look at panel data, which comprises both of a cross-sectional and time-series variation. Panel data is quite useful because if we had omitted **time-invariant** variables (i.e. omitted variables that do not change over time), then we can overcome this omitted variable problem by allowing for individual-specific intercepts which will control for these omitted variables. That is, the intercept will absorb the omitted variable error. We will be looking small T panel datasets and taking number of cross-sectional observations $n \rightarrow \infty$ when we do asymptotics.

13.1 Two-Way Error Components Model

We will test and impose restrictions on our panel dataset that require some of the estimated parameters to be common to all the observation units. These restrictions are known as **pooling restrictions**, and estimators which impose these restrictions are known as **pooled estimators**.

We will have three different types of variables.

1. Observations that vary both across individual observation unit i and over time t . An example of this could be someone's wage over their lifetime. Each person will have a different wage, and for a given person, their wage will vary over the course of their lifetime.
2. Observations that vary across individual observations i but not over time. An example of this could be someone's name or gender whereby different people will have different names but for that person, their name or gender will remain the same throughout the course of their life.
3. Observations that vary across time t but not individual observations. An example of this could be the interest rate on government bonds that individuals can save in. Everyone has access to the same interest rate, so this does not vary across individuals. However, the interest rate on government bonds does vary over time.

As we have three different types of variables, we can also define three different types of errors associated to these three different variables.

Definition 561 (Panel Data Explanatory Variables and Errors) 1. $\mathbf{x}_{it}$: Explanatory variables that vary over individuals and time.

2. $\mathbf{w}_i$: Explanatory variables that vary over individuals.

3. $\mathbf{s}_t$: Explanatory variables that vary over time.

The associated error terms are:

1. v_{it} for variables that vary over individuals and time,

2. η_i for variables that vary over individuals,

3. f_t for variables that vary over time.

Remark 562. The best way to memorise this is to look at the subscript of the variable. If the variable has both i for individual and t for time, so in this $\mathbf{x}_{it}$, then it is a variable that is varying across time and individuals.

If the variable only has the t subscript, such as $\mathbf{s}_t$, then we automatically know it must be the explanatory variable that is varying over time but is the same across individuals.

We can now state the linear model involving these variables.

Definition 563 (Two-Way Error Components Model)

The two-way error components model is:

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \mathbf{w}_i\boldsymbol{\gamma} + \mathbf{s}_t\boldsymbol{\delta} + (v_{it} + \eta_i + f_t). \quad (485)$$

Remark 564. *It is called two-way error even though we have three error terms!*

Remark 565. *As we will later see, we have not said anything about η which is known as the unobserved heterogeneity between individuals. In particular, we have not said anything about its correlation with $\mathbf{X}_i$. This is quite important and will lead to assumptions known as fixed and random effects later on.*

13.1.1 Time Dummies

Before we proceed, we will now show how to deal with the time variable $\mathbf{s}_t$ and associated error term f_t in the two-way error components model quite easily. The idea is that if we include T dummy variables, which acts as an intercept, we can absorb the effect of these time-varying variables into these T different intercepts. We now formally show this.

First, start off with the two-way error components model:

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \mathbf{w}_i\boldsymbol{\gamma} + \mathbf{s}_t\boldsymbol{\delta} + (v_{it} + \eta_i + f_t).$$

Now, we want to combine the time-varying explanatory variable $\mathbf{s}_t$ and the time-varying error component f_t as:

$$\phi_t = \mathbf{s}_t\boldsymbol{\delta} + f_t. \quad (486)$$

We can then substitute this into the two-way error components model to get:

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \mathbf{w}_i\boldsymbol{\gamma} + (\mathbf{s}_t\boldsymbol{\delta} + f_t) + (v_{it} + \eta_i).$$

which then becomes:

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \mathbf{w}_i\boldsymbol{\gamma} + \phi_t + (v_{it} + \eta_i).$$

Therefore, we can absorb this into the error term and we now have the new equation:

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \mathbf{w}_i\boldsymbol{\gamma} + (v_{it} + \eta_i + \phi_t). \quad (487)$$

We now define time dummy variables.

Definition 566 (Time Dummies)

Time dummy variables are dummy variables which indicate which period of time we are in:

$$D_t^j = \begin{cases} 1 & \text{for observations in period } j \\ 0 & \text{otherwise.} \end{cases} \quad (488)$$

for $j = 1, \dots, T$.

Therefore, we can take our equation:

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \mathbf{w}_i\boldsymbol{\gamma} + (v_{it} + \eta_i + \phi_t).$$

and notice that

$$D_t^s \phi_t = \phi_t \quad \text{if } s = t.$$

Therefore, we can therefore rewrite our earlier equation as:

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \mathbf{w}_i\boldsymbol{\gamma} + \phi_1 D_t^1 + \phi_2 D_t^2 + \dots + \phi_T D_t^T + (v_{it} + \eta_i).$$

We do not have an intercept in $\mathbf{x}_{it}$ and therefore we do not suffer from a **dummy variable trap** when we include these T dummy variables. This equation has period-specific intercepts whereby for each period s, we will have a different intercept. This intercept will then absorb the error term f_t . This equation has a special name.

Definition 567 (One-way Error Components Model)

The one-way error components model is given by:

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \mathbf{w}_i\boldsymbol{\gamma} + \sum_{s=1}^T \phi_s D_t^s + (v_{it} + \eta_i) \quad (489)$$

where $\mathbf{x}_{it}$ is a $1 \times K$ vector, $\boldsymbol{\beta}$ is a $K \times 1$ vector, $\mathbf{w}_i$ is a $1 \times G$ vector and $\boldsymbol{\gamma}$ is a $G \times 1$ vector.

With the one-way error components model, we can estimate the T parameters $(\phi_1, \dots, \phi_T)$ alongside the original $\boldsymbol{\beta}$ and $\boldsymbol{\gamma}$ vectors. With small T panels, losing identification of the δ parameter associated to the time varying variable s_t is not a concern. We also say that this model is **saturated** in the time dimension.

Writing the one-way error components model is quite a useful tool as this means we can now not worry about time-varying $s_t\delta + f_t$ terms as we can simply include these time dummy variables into our model and take care of them.

13.2 Pooled OLS in Linear Panel Model

From our discussion in the last section, we can focus on the one-way error components model whereby we take care of the time varying variables by simply including a dummy variable. Hence, we are interested in the model:

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \mathbf{w}_i\boldsymbol{\gamma} + (v_{it} + \eta_i)$$

for $i = 1, \dots, N$ and $t = 1, \dots, T$. Now, we want to stack the T observations for each individual observation i to get:

$$\mathbf{y}_i = \mathbf{X}_i\boldsymbol{\beta} + \mathbf{W}_i\boldsymbol{\gamma} + (\mathbf{v}_i + \eta_i \mathbf{j}_T)$$

whereby the dimensions are:

1. $\mathbf{y}_i$ is a $T \times 1$ vector which represents all T time observations for individual i.
2. $\mathbf{X}_i$ is a $T \times K$ matrix where each row t is the K explanatory variables for observation i at time t and each column j represents all T time observations for individual i for explanatory variable j.
3. $\mathbf{W}_i$ is a $T \times G$ matrix of time-invariant variables. Each row contains the G time-invariant variables. Each row is identical as these variables do not vary over time.
4. $\mathbf{v}_i$ is a $T \times 1$ vector where each row t contains the error term for observation i at time t.
5. $\eta_i \mathbf{j}_T$ is a $T \times 1$ vector where η_i is a scalar and $\mathbf{j}_T$ is a $T \times 1$ vector full of ones.

Now, we can stack over the N individuals to get:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{W}\boldsymbol{\gamma} + (\boldsymbol{\eta} + \mathbf{v}) \quad (490)$$

where the dimensions are:

1. $\mathbf{y}$ is a $NT \times 1$ vector.
2. $\mathbf{X}$ is a $NT \times K$ matrix.
3. $\mathbf{W}$ is a $NT \times G$ matrix.
4. $\mathbf{v}$ is a $NT \times 1$ vector.
5. $\boldsymbol{\eta}$ is a $NT \times 1$ vector.

Before, we proceed, we make an assumption.

Assumption 568 (No Time-Invariant Variables)

We assume that there are no time-invariant explanatory variables.

$$\mathbf{W} = \mathbf{0}. \quad (491)$$

However, note that even if $\mathbf{W} = \mathbf{0}$, we can still have time invariant errors $\boldsymbol{\eta}_i$. Therefore, our model is now:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + (\boldsymbol{\eta} + \mathbf{v}).$$

Now, if we want to carry out estimation, we need **two** more assumptions over:

$$y_i = \mathbf{X}_i\boldsymbol{\beta} + (\boldsymbol{\eta}_i j_T + \mathbf{v}_i).$$

Assumption 569 (Cross-Sectional Independence)

Observations on $(y_i, \mathbf{X}_i)$ are independent over $i = 1, \dots, N$.

Remark 570. This says that if you know $(\mathbf{X}_i, y_i)$ for one person, it tells you nothing about person j 's $(\mathbf{X}_j, y_j)$.

Assumption 571 (Slope Parameter Homogeneity)

The parameters in $\boldsymbol{\beta}$ are common to all observations $i = 1, \dots, N$.

We can now estimate $\boldsymbol{\beta}$ under all these assumptions.

Definition 572 (Pooled OLS Estimator)

Suppose we had no time-invariant variables $\mathbf{w}_i$. We then have the structural linear panel model:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u} \quad (492)$$

where $\mathbf{u} = \boldsymbol{\eta} + \mathbf{v}$. Furthermore, assume that we had cross-sectional independence and slope parameter homogeneity. Then, the **pooled ordinary least squares** estimator is:

$$\hat{\boldsymbol{\beta}}_{OLS} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}. \quad (493)$$

Remark 573. *The pooled OLS is also known as the **levels OLS**.*

The pooled OLS corresponds to running OLS on the observations pooled across individuals i over time t and then computing the OLS.

We now want to establish asymptotic properties of the pooled OLS estimator. To do so, we require more assumptions. We now go through them.

Assumption 574 (Errors Have Expected Value of Zero)

The error components have expected value of zero:

$$\mathbb{E}[u_{it}] = \mathbb{E}[\eta_i] + \mathbb{E}[v_{it}] = 0. \quad (494)$$

Remark 575. *This assumption is not actually not needed but this makes notation easier since we do not need to include intercept terms in our analysis.*

Recall that when we looked at consistency in the setting of OLS, we require assumptions on orthogonality between the regressor $\mathbf{X}$ and the error term $\mathbf{u}$. In panel data, we need to do a similar thing whereby we need to impose orthogonality between the explanatory variables $\mathbf{x}_{it}$ and the error terms v_{it} and η_i . In fact, for this entire section on panel data, the aim is to know which orthogonality conditions we will need between $\mathbf{x}_{it}$ and the error terms. Hence, pay close attention to each time we mention what is the orthogonality condition that we are imposing.

We now discuss what assumptions we shall impose on the correlation between the error term $u_{it} = \eta_i + v_{it}$ and $\mathbf{x}_{it}$.

Assumption 576 (Contemporaneous Exogeneity)

The K explanatory variables are uncorrelated with the time varying individual error term in the same time period:

$$\mathbb{E}[\mathbf{x}_{it} v_{it}] = 0 \quad (495)$$

for $t = 1, 2, \dots, T$.

Remark 577. *This says that $\mathbf{x}_{it}$ is uncorrelated with the individual time-varying error term v_{it} in every period t .*

Note that under contemporaneous exogeneity, it may still be the case that:

$$\mathbb{E}[\mathbf{x}_{it} v_{i,t-k}] \neq 0$$

for $k > 0$ so that $\mathbf{x}_{it}$ may be correlated with lagged errors $v_{i,t-k}$.

We now need to discuss the correlation between $\mathbf{x}_{it}$ and the time-invariant (individual specific) error term η_i .

Assumption 578 (Uncorrelated Individual Effects/Random Effects)

All K explanatory variables $\mathbf{x}_{it}$ is uncorrelated with all individual-specific effects η_i :

$$\mathbb{E}[\mathbf{x}_{it} \eta_i] = 0 \quad (496)$$

for $t = 1, 2, \dots, T$.

We have stated the two assumptions we need for orthogonality between $\mathbf{x}_{it}$ and the two error terms. Before we proceed, we need one more assumption that we also needed in OLS.

Assumption 579 (Full Rank)

The matrix $\mathbf{X}$ is full rank.

We now collect all the assumptions together to recall what we have assumed so far.

Assumption 580 (Assumptions for Consistency of Pooled OLS)

1. $(y_i, \mathbf{X}_i)$ are independent over $i = 1, \dots, N$.
2. $\boldsymbol{\beta}$ are common to all observations $i = 1, \dots, N$.
3. The matrix $\mathbf{X}$ is full rank.
4. $\mathbb{E}[\mathbf{x}_{it}\eta_i] = \mathbf{0}$ for $t = 1, 2, \dots, T$.
5. $\mathbb{E}[\mathbf{x}_{it}v_{it}] = \mathbf{0}$ for $t = 1, 2, \dots, T$.
6. $\mathbb{E}[u_{it}] = \mathbb{E}[\eta_i] + \mathbb{E}[v_{it}] = 0$.

Recall that the linear panel model is:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}$$

where $\mathbf{u} = \boldsymbol{\eta} + \mathbf{v}$. Under these assumptions for consistency of pooled OLS, we have that:

$$\mathbb{E}[\mathbf{x}_{it}u_{it}] = \mathbb{E}[\mathbf{x}_{it}\eta_i] + \mathbb{E}[\mathbf{x}_{it}v_{it}] = \mathbf{0}.$$

This is the **orthogonality condition** that we have previously seen in linear regression that we needed in order to establish consistency.

Proposition 581 (Consistency of Pooled OLS Estimator)

Suppose we have the linear panel model:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u} \tag{497}$$

where $\mathbf{u} = \boldsymbol{\eta} + \mathbf{v}$. Then, the assumptions for consistency of pooled OLS, we have that the pooled OLS estimator is a consistent estimator:

$$\hat{\boldsymbol{\beta}}_{OLS} \xrightarrow{P} \boldsymbol{\beta} \tag{498}$$

as $N \rightarrow \infty$ with T fixed. Joint (N, T) asymptotics require additional time-series dependence conditions.

We have shown consistency and now want to look at inference. If we wanted to compute standard errors for inference for the pooled OLS, we notice that we have some minor issues as the error terms $u_{it} = \eta_i + v_{it}$ are **serially correlated**.

Proposition 582 (Serial Correlation of Error Terms in Pooled OLS)

Under the standard error-components assumptions with $\text{Var}(\eta_i) > 0$ and serially uncorrelated v_{it} , the composite errors $u_{it} = \eta_i + v_{it}$ are **serially correlated**. As a result, ordinary pooled OLS is generally **not efficient**.

Proof. If v_{it} is serially uncorrelated and independent of η_i , the errors u_{it} are correlated through the time-invariant η_i

term as:

$$u_{i,t} = \eta_i + v_{i,t}$$

$$u_{i,t-1} = \eta_i + v_{i,t-1}$$

η_i appears in both terms, and $Cov(u_{it}, u_{is}) = Var(\eta_i) > 0$ for $t \neq s$. This violates the spherical-error condition used for OLS efficiency, so pooled OLS is generally not efficient. $\square$

As a result of this serial correlation, we need to use **cluster-robust standard errors and test statistics** whereby we treat all T observations for individual i as a cluster to help us rectify this source of serial correlation.

13.3 Fixed Effects Model

First, we recall from last section the uncorrelated individual effects assumption we made:

$$\mathbb{E}[\mathbf{x}_{it}\eta_i] = \mathbf{0}$$

for $t = 1, 2, \dots, T$. This is in fact a **strong and restrictive assumption**. We can in fact relax this assumption whereby we now allow some explanatory variables $\mathbf{x}_{it}$ to be correlated with the individual-specific effects η_i .

We first describe the setting we are in.

Assumption 583 (Short T Panel)

The panel set has a small number of time observations T .

Now, we state the new relaxed assumption we will make.

Assumption 584 (Correlated Individual Effects/Fixed Effects)

Some of the K explanatory variables $\mathbf{x}_{it}$ are correlated with the individual-specific effects η_i :

$$\mathbb{E}[\mathbf{x}_{it}\eta_i] \neq \mathbf{0} \quad (499)$$

for $t = 1, 2, \dots, T$.

For the other error term v_{it} , we keep the same assumption we made in pooled OLS.

Assumption 585 (Contemporaneous Exogeneity)

The K explanatory variables are uncorrelated with the time varying individual error term in the same time period:

$$\mathbb{E}[\mathbf{x}_{it}v_{it}] = \mathbf{0} \quad (500)$$

for $t = 1, 2, \dots, T$.

Remark 586. *This rules out contemporaneous correlation between $\mathbf{x}_{it}$ and v_{it} .*

Recalling the section on instrumental variables, this η_i can be thought of as containing omitted variables and hence capturing the omitted variables in our model. As a result of this condition, if we go back to the linear panel model:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{u}$$

where $\mathbf{u} = \boldsymbol{\eta} + \mathbf{v}$, we can think of $\boldsymbol{\eta}$ as capturing omitted variables. To not have to worry about adding intercepts, we assume that the error components have expected values of zero:

$$\mathbb{E}[\eta_i] = \mathbb{E}[v_{it}] = 0.$$

Under this new assumption of correlated individual effects, we no longer have orthogonality assumption between $\mathbf{x}_{it}$ and the error terms:

$$\mathbb{E}[\mathbf{x}_{it}u_{it}] = \mathbb{E}[\mathbf{x}_{it}\eta_i] + \mathbb{E}[\mathbf{x}_{it}v_{it}] = \mathbb{E}[\mathbf{x}_{it}\eta_i] + 0 \neq \mathbf{0}.$$

Therefore, we no longer have the orthogonality condition $\mathbb{E}[\mathbf{x}_{it}u_{it}] = 0$ needed for pooled OLS to be consistent. As a result of this, we need to think of other ways to estimate $\boldsymbol{\beta}$. What we shall do is to transform the original model

using the fact that we now have repeated observations over time. This will eliminate the time-invariant components from the error term of the transformed equation. This will then allow us to derive consistent estimators of β by using OLS on the transformed model. The issue though from this approach is that even though we eliminated η_i from the error term, we now require new exogeneity assumptions instead of the predetermined assumption which is what we used in the uncorrelated individual effects model, which mathematically was $\mathbb{E}[\mathbf{x}_{it}v_{it}] = \mathbf{0}$.

We will now look at a few different estimators for β under the assumption of correlated individual effects.

13.3.1 Within Groups Estimator

We look at the first transformation we can use. The idea is to take the average of each variable over its time component and demean the original variable by its time average. Intuitively, as the nuisance time-invariant error term η_i does not vary over time, this transformation will eliminate it. We now formally state this transformation.

Definition 587 (Within Transformation)

The within transformation is given by:

$$\tilde{y}_{it} = y_{it} - \bar{y}_i \quad (501)$$

$$\tilde{\mathbf{x}}_{it} = \mathbf{x}_{it} - \bar{\mathbf{x}}_i \quad (502)$$

$$\tilde{\eta}_i = \eta_i - \bar{\eta}_i \quad (503)$$

$$\tilde{v}_{it} = v_{it} - \bar{v}_i \quad (504)$$

where the means are defined as:

$$\bar{y}_i = \frac{1}{T} \sum_{s=1}^T y_{is}$$

$$\bar{\mathbf{x}}_i = \frac{1}{T} \sum_{s=1}^T \mathbf{x}_{is}$$

$$\bar{\eta}_i = \frac{1}{T} \sum_{s=1}^T \eta_i$$

$$\bar{v}_i = \frac{1}{T} \sum_{s=1}^T v_{is}.$$

Under this within transformation, something interesting happens to the error terms.

Proposition 588 (Within Transformation of Error Terms)

The within transformation of the time-invariant error term η_i is zero:

$$\tilde{\eta}_i = 0. \quad (505)$$

Furthermore, for individual i 's time-varying error term v_{it} , its within transformation depends on all of its time-varying error terms:

$$\tilde{v}_{it} = v_{it} - \frac{1}{T} (v_{i1} + v_{i2} + \dots + v_{iT}) \quad (506)$$

whereby $v_{i1}, \dots, v_{iT}$ are all the time varying error terms for individual i .

Proof. For the time-invariant error term's within transformation:

$$\tilde{\eta}_i = \eta_i - \frac{1}{T} \sum_{t=1}^T \eta_i = \eta_i - \frac{1}{T} (T \eta_i) = \eta_i - \eta_i = 0.$$

For the time-varying error term's within transformation:

$$\tilde{v}_{it} = v_{it} - \bar{v}_i = v_{it} - \frac{1}{T} (v_{i1} + v_{i2} + \dots + v_{iT}).$$

□

Therefore, we have just shown that the time-invariant error term η_i has been eliminated under the within transformation. This is good news because under our correlated individual effects assumption, we had that:

$$\mathbb{E}[\mathbf{x}_{it}\eta_i] \neq \mathbf{0}$$

which prevented us from having consistency in our pooled OLS estimator. We have now successfully removed this error term η_i from our model and therefore don't have to worry about endogeneity between our error term η_i and regressor $\mathbf{x}_{it}$.

However, an issue that has arisen now is that $\tilde{v}_{it}$ depends on **all** of its time-varying components $v_{i1}, v_{i2}, \dots, v_{iT}$. This matters for our previous assumption of exogeneity:

$$\mathbb{E}[\mathbf{x}_{it}v_{it}] = \mathbf{0}$$

will **not** guarantee for us that:

$$\mathbb{E}[\tilde{\mathbf{x}}_{it}\tilde{v}_{it}] = \mathbf{0}.$$

To be more explicit on why this is the case, recall that the explicit expressions of the within transformations of the explanatory variable and time-varying error terms are:

$$\begin{aligned}\tilde{\mathbf{x}}_{it} &= \mathbf{x}_{it} - \frac{1}{T} (\mathbf{x}_{i1} + \mathbf{x}_{i2} + \dots + \mathbf{x}_{iT}) \\ \tilde{v}_{it} &= v_{it} - \frac{1}{T} (v_{i1} + v_{i2} + \dots + v_{iT}).\end{aligned}$$

Therefore, when we try to impose an exogeneity assumption on this:

$$\mathbb{E}[\tilde{\mathbf{x}}_{it}\tilde{v}_{it}] = \mathbb{E}\left[\left(\mathbf{x}_{it} - \frac{1}{T} (\mathbf{x}_{i1} + \mathbf{x}_{i2} + \dots + \mathbf{x}_{iT})\right)\left(v_{it} - \frac{1}{T} (v_{i1} + v_{i2} + \dots + v_{iT})\right)\right]$$

we can see that whilst our previous exogeneity assumption of $\mathbb{E}[\mathbf{x}_{it}v_{it}] = 0$ will eliminate some of the expectation of the products, there will be many terms in the above expression where:

$$\mathbb{E}[\mathbf{x}_{it}v_{is}] \neq \mathbf{0}$$

for $s \neq t$ which means we do not have exogeneity for OLS. Therefore, we need to impose a new assumption called strict exogeneity.

Assumption 589 (Strict Exogeneity)

All K explanatory variables are uncorrelated with individual time-varying error terms over all periods:

$$\mathbb{E}[\mathbf{x}_{it}v_{is}] = \mathbf{0} \quad (507)$$

for all $s, t \in \{1, 2, \dots, T\}$.

Strict exogeneity requires that the current values of $\mathbf{x}_{it}$ are uncorrelated with past, present, and future values of the time-varying component of the error term v_{is} for the same individual in the original model.

An issue with the assumption of strict exogeneity is that it rules out allowing lagged dependent variables $y_{i,t-1}$ as an explanatory variable in $\mathbf{x}_{it}$, which is what we have previously done for autoregressive models in time series. The reason we can't do so is because $y_{i,t-1}$ depends on previous error terms $v_{i,t-1}$ by definition of the structural equation. Therefore, we will come back later on in the course to learn how to allow for lagged dependent variables as an explanatory variable.

We can now estimate β .

Definition 590 (Within Groups/Fixed Effects Estimator)

Suppose that our structural model is:

$$\tilde{y}_{it} = \tilde{\mathbf{x}}_{it}\beta + \tilde{v}_{it}. \quad (508)$$

where $\tilde{y}_{it}, \tilde{\mathbf{x}}_{it}, \tilde{v}_{it}$ are the within-transformations. Then, stack the observations to form the structural equation:

$$\tilde{\mathbf{y}} = \tilde{\mathbf{X}}\beta + \tilde{\mathbf{v}}. \quad (509)$$

The **within groups estimator** of β is given by:

$$\hat{\beta}_{WG} = \left(\tilde{\mathbf{X}}^T \tilde{\mathbf{X}} \right)^{-1} \tilde{\mathbf{X}}^T \tilde{\mathbf{y}}. \quad (510)$$

Remark 591. The within groups estimator is also known as the **fixed effects estimator**.

We also have all the necessary orthogonality conditions needed for consistency of the within groups estimator.

Proposition 592 (Consistency of Within Groups Estimator)

Under the assumption of strict exogeneity, the within groups estimator is a consistent estimator:

$$\hat{\beta}_{WG} \xrightarrow{P} \beta \quad (511)$$

as $N \rightarrow \infty$ with T fixed.

It is important to note that we have just shown that $\hat{\beta}_{WG}$ is a consistent estimator for β even if all explanatory variables $\mathbf{x}_{it}$ are correlated with η_i :

$$\mathbb{E}[\mathbf{x}_{it}\eta_i] \neq \mathbf{0}$$

which is what we called the correlated individuals effects/fixed effects. We can also derive the asymptotic distribution. First, we need an assumption.

Assumption 593 (Panel Data Homoskedasticity)

Panel data homoskedasticity assumes that the time varying individual error terms v_{it} satisfy:

$$v_{it}|\mathbf{X} \sim i.i.d(0, \sigma_v^2) \quad (512)$$

for all observations i and time periods t .

We can now derive the asymptotic distribution of the estimator.

Theorem 594 (Asymptotic Normality of Within Groups Estimator)

Under strict exogeneity and panel data homoskedasticity, the asymptotic distribution of the within groups estimator is:

$$\sqrt{N}(\hat{\boldsymbol{\beta}}_{WG} - \boldsymbol{\beta}) \xrightarrow{D} \mathcal{N}(\mathbf{0}, \sigma_v^2 \mathbf{Q}^{-1}), \quad (513)$$

as $N \rightarrow \infty$ with T fixed, where $\mathbf{Q} = \lim_{N \rightarrow \infty} N^{-1} \tilde{\mathbf{X}}^T \tilde{\mathbf{X}}$. A consistent estimator $\hat{\sigma}_v^2$ can be used in the corresponding estimated covariance matrix.

We will discuss the consistent estimator term $\hat{\sigma}_v^2$ soon.

We note that even if we don't have homoskedasticity, this **does not** affect the consistency of the within groups estimator as consistency only requires strict exogeneity from our earlier arguments.

Furthermore, if panel data homoskedasticity does not hold, we can use heteroskedasticity-consistent and cluster-robust standard errors. As a result, we do not need to worry too much about serial correlation in the error terms.

Recall that in the asymptotic distribution of the within groups estimator, we had a $\hat{\sigma}_v^2$ estimator. We now describe what we can use as a consistent estimator $\hat{\sigma}_v^2$.

Proposition 595 (Consistent Estimator of the Variance)

Under the panel data homoskedasticity assumption, we can estimate the variance of v_{it} , denoted by σ_v^2 , by constructing the estimator:

$$\hat{\sigma}_v^2 = \frac{\sum_{i=1}^N \sum_{t=1}^T \hat{v}_{it}^2}{N(T-1) - K} = \frac{\hat{\mathbf{v}}^T \hat{\mathbf{v}}}{N(T-1) - K} \quad (514)$$

where $\hat{v}_{it}$ are the residuals from the within groups estimator:

$$\hat{v}_{it} = \tilde{y}_{it} - \tilde{\mathbf{x}}_{it} \hat{\boldsymbol{\beta}}_{WG}. \quad (515)$$

Remark 596. We will later see where the $N(T-1) - K$ degrees of freedom come from when we look at least squares dummy variables estimators.

13.3.2 Least Squares Dummy Variables Estimator

This is another estimator we can use to estimate $\boldsymbol{\beta}$ under the assumption of fixed effects. If we do not have too many cross-sectional observations n , we can include N dummy variables in order to absorb the error term η_i which we are having issues with. This is a similar idea to what we did for the time dummy variables.

Definition 597 (Dummy Variables)

We define the N dummy variables for the N observations as:

$$I_i^1 = 1 \quad \text{for observations on individual 1 or otherwise } I_i^1 = 0 \quad (516)$$

$$I_i^2 = 1 \quad \text{for observations on individual 2 or otherwise } I_i^2 = 0 \quad (517)$$

$$\dots \quad (518)$$

$$I_i^N = 1 \quad \text{for observations on individual } N \text{ or otherwise } I_i^N = 0 \quad (519)$$

With these N dummy variables, we now include them into our structural model and describe our new estimation strategy.

Definition 598 (Least Squares Dummy Variables Estimator)

Suppose that the structural model with dummy variables is now:

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \eta_1 I_i^1 + \eta_2 I_i^2 + \dots + \eta_N I_i^N + v_{it}. \quad (520)$$

Let $\mathbf{D}$ collect the individual dummy variables and let $\mathbf{M}_D = \mathbf{I}_{NT} - \mathbf{D}(\mathbf{D}^T \mathbf{D})^{-1} \mathbf{D}^T$. The **Least Squares Dummy Variables Estimator** (LSDV) for the slope coefficients is the OLS estimator in the above structural equation:

$$\hat{\boldsymbol{\beta}}_{LSDV} = (\mathbf{X}^T \mathbf{M}_D \mathbf{X})^{-1} \mathbf{X}^T \mathbf{M}_D \mathbf{y}. \quad (521)$$

The LSDV estimator should produce the same coefficients as the fixed effect (FE) estimator.

Theorem 599 (Equivalency of LSDV and WG)

The LSDV estimator and WG estimator are equivalent:

$$\hat{\boldsymbol{\beta}}_{LSDV} = \hat{\boldsymbol{\beta}}_{WG}. \quad (522)$$

When based on the same error-variance estimator and degrees-of-freedom adjustment, the two methods also produce the same conventional standard errors for $\boldsymbol{\beta}$. Their main difference is computational: the within transformation avoids estimating a large set of individual dummy coefficients when N is large.

If we look at the LSDV structural equation, we have our usual K parameters we estimate in $\boldsymbol{\beta}$. We also have an additional N parameters $\eta_1, \dots, \eta_N$ which we estimate. Therefore, the degrees of freedom is of our estimation is:

$$NT - K - N = N(T - 1) - K$$

which is what we saw earlier in Eq. (514) when computing the consistent estimator of the variance under panel data homoskedasticity assumption.

13.3.3 First-Differenced OLS

This is another estimator we can use to estimate $\boldsymbol{\beta}$ under the fixed effects assumptions. We can perform another transformation that is inspired by time series techniques.

Definition 600 (First-Differencing)

The first differencing operator Δ is defined as the operator such that when applied to y_{it} results in:

$$\Delta y_{it} = y_{it} - y_{i,t-1}. \quad (523)$$

Now, we can take our original structural model:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\eta} + \mathbf{v}$$

and apply our first-difference transformation on it.

Proposition 601 (First-Difference Transformation)

The first-differenced transformed equation of our panel data structural model is:

$$\Delta y_{it} = \Delta \mathbf{x}_{it}\boldsymbol{\beta} + \Delta v_{it} \quad (524)$$

for $i = 1, 2, \dots, N$ and $t = 2, 3, \dots, T$.

Proof. Start off with the panel data structural model:

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \eta_i + v_{it}$$

First, we can simply difference y_{it} to get:

$$y_{it} - y_{i,t-1} = \Delta y_{it}$$

and the same applies for $\mathbf{x}_{it}, v_{it}$:

$$\mathbf{x}_{it} - \mathbf{x}_{i,t-1} = \Delta \mathbf{x}_{it}.$$

$$v_{it} - v_{i,t-1} = \Delta v_{it}.$$

However, notice that for the time-invariant parameter η_i , differencing it removes it:

$$\eta_i - \eta_i = 0.$$

□

Notice that as a result of first-differencing, we start from $t = 2$ as we do not observe Δy_{i1} and $\Delta \mathbf{x}_{i1}$. We can now state our estimator under this transformation.

Definition 602 (First-Differenced OLS)

Suppose that the structural model under first-differencing is:

$$\Delta y_{it} = \Delta \mathbf{x}_{it}\boldsymbol{\beta} + \Delta v_{it} \quad (525)$$

for $i = 1, 2, \dots, N$ and $t = 2, 3, \dots, T$.

Then, the first-differenced OLS estimator is:

$$\hat{\boldsymbol{\beta}}_{FDOLS} = [(\Delta \mathbf{X})^T \Delta \mathbf{X}]^{-1} (\Delta \mathbf{X})^T (\Delta \mathbf{y}). \quad (526)$$

If we look at the first-differenced structural equation:

$$\Delta y_{it} = \Delta \mathbf{x}_{it} \boldsymbol{\beta} + \Delta v_{it}$$

we can see that in order for usual orthogonality condition for consistency of the OLS estimator, we require that:

$$\mathbb{E}[\Delta \mathbf{x}_{it} \Delta v_{it}] = \mathbf{0}.$$

Assumption 603 (First Differenced Exogeneity)

The first differenced of $\mathbf{x}_{it}$ and the first differenced time varying individual error v_{it} are uncorrelated:

$$\mathbb{E}[\Delta \mathbf{x}_{it} \Delta v_{it}] = \mathbf{0}. \quad (527)$$

We can now state our result for consistency of this estimator.

Proposition 604 (Consistency of First-Differenced OLS)

Under the assumption of first differenced exogeneity, the first-differenced OLS estimator is consistent:

$$\hat{\boldsymbol{\beta}}_{FDOLS} \xrightarrow{P} \boldsymbol{\beta} \quad (528)$$

as $N \rightarrow \infty$ with T fixed.

First-differenced exogeneity restricts a combination of contemporaneous and adjacent-period cross-moments. By itself, it does not set each of those cross-moments to zero separately. We can nevertheless establish one clear relationship with strict exogeneity.

Proposition 605 (Strict Exogeneity Implies First Differenced Exogeneity)

Strict exogeneity

$$\mathbb{E}[\mathbf{x}_{it} v_{is}] = \mathbf{0} \quad (529)$$

for $t, s \in \{1, \dots, T\}$ implies first differenced exogeneity:

$$\mathbb{E}[\Delta \mathbf{x}_{it} \Delta v_{it}] = \mathbf{0} \quad (530)$$

for $t = 1, \dots, T$.

Proof. We directly show this by definition of first differenced exogeneity:

$$\begin{aligned} \mathbb{E}[\Delta \mathbf{x}_{it} \Delta v_{it}] &= \mathbb{E}[(\mathbf{x}_{it} - \mathbf{x}_{i,t-1})(v_{it} - v_{i,t-1})] \\ &= \mathbb{E}[\mathbf{x}_{it} v_{it}] - \mathbb{E}[\mathbf{x}_{it} v_{i,t-1}] - \mathbb{E}[\mathbf{x}_{i,t-1} v_{it}] + \mathbb{E}[\mathbf{x}_{i,t-1} v_{i,t-1}] \\ &= \mathbf{0} + \mathbf{0} + \mathbf{0} + \mathbf{0} \end{aligned}$$

where the last equation holds by strict exogeneity. □

Remark 606. First-differenced exogeneity does not generally imply contemporaneous exogeneity $\mathbb{E}[\mathbf{x}_{it} v_{it}] = \mathbf{0}$, nor does contemporaneous exogeneity generally imply first-differenced exogeneity. Additional restrictions are needed to order these two assumptions. Strict exogeneity implies both.

Therefore, this means that under the assumptions of strict exogeneity, then the first differenced OLS estimator is consistent!

We can also state what the standard errors are for FDOLS.

Theorem 607 (Standard Errors for First Differenced OLS Estimator)

Classical standard errors are appropriate for $\hat{\beta}_{FDOLS}$ if the time-varying error component v_{it} is a random walk:

$$v_{it} = v_{i,t-1} + \epsilon_{it} \quad (531)$$

with i.i.d innovations so that the first differenced of the error terms:

$$\Delta v_{it} = \epsilon_{it} \sim i.i.d(0, \sigma_\epsilon^2) \quad (532)$$

for all individuals i and time t .

Again, consistency for $\hat{\beta}_{FDOLS}$ **does not** rely on v_{it} being uncorrelated but instead relies on first differenced exogeneity.

13.3.4 Efficiency Comparison of Fixed Effects Estimators

We restrict our attention to a special case of endogeneity.

Assumption 608 (Correlation with Time-Invariant Error Terms)

All K explanatory variables are correlated with η_i

Recall that efficient estimators are estimators that achieves the minimum possible variance as an estimator. We now discuss the conditions for efficiency for the estimators we have seen so far.

Theorem 609 (Efficiency of Within Groups Estimator)

Suppose that we have the classical assumptions:

$$v_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT} \sim i.i.d(0, \sigma_v^2) \quad (533)$$

over $i = 1, \dots, N$ and $t = 1, 2, \dots, T$. Then, $\hat{\beta}_{WG}$ is an efficient estimator of β . Furthermore, under normality assumption:

$$v_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT} \sim i.i.d\mathcal{N}(0, \sigma_v^2) \quad (534)$$

then $\hat{\beta}_{WG}$ coincides with the maximum likelihood estimator of β and achieves the Cramer-Rao lower bound.

Theorem 610 (Efficiency of First Differenced Estimator)

Suppose that we have that

$$v_{it} = v_{i,t-1} + \epsilon_{it} \quad (535)$$

is a random walk and the innovations

$$\epsilon_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT} \sim i.i.d(0, \sigma_\epsilon^2) \quad (536)$$

over $i = 1, \dots, N$ and $t = 1, 2, \dots, T$. Then, $\hat{\beta}_{FDOLS}$ is an efficient estimator of β .

We now compare these two estimators and describe when is one more efficient than the other.

Proposition 611 (Comparison of the Two Fixed Effects Estimators)

Suppose that v_{it} is serially uncorrelated, then $\hat{\beta}_{WG}$ is more efficient than $\hat{\beta}_{FDOLS}$.

Suppose that v_{it} is a random walk, then $\hat{\beta}_{FDOLS}$ is more efficient than $\hat{\beta}_{WG}$.

If we are not worried about serial correlation in v_{it} and there is smooth continuous time series variation in y_{it} and x_{it} , then either estimator is fine.

Remark 612. *These two estimators coincide when we only have $T = 2$ time periods in our panel dataset.*

A big issue to worry about for first-differenced OLS is that it may not do well when you have step functions in your explanatory variables.

13.3.5 Large T Asymptotics of Fixed Effects Estimators

In this section, we explain how strict exogeneity can be relaxed under large- T asymptotics, subject to suitable weak-dependence and moment conditions. Recall strict exogeneity:

$$\mathbb{E}[x_{it}v_{is}] = 0$$

for $t, s \in \{1, \dots, T\}$ for the WG estimator but not for the FDOLS estimator.

First, recall that for the within group transformation:

$$\begin{aligned}\tilde{y}_{it} &= y_{it} - \frac{1}{T} \sum_{s=1}^T y_{is} \\ \tilde{x}_{it} &= x_{it} - \frac{1}{T} \sum_{s=1}^T x_{is} \\ \tilde{v}_{it} &= v_{it} - \frac{1}{T} \sum_{s=1}^T v_{is}.\end{aligned}$$

As $T \rightarrow \infty$, the time averages do not literally disappear. Rather, under suitable weak-dependence and ergodicity conditions, the endogeneity induced by demeaning can become asymptotically negligible. A standard weaker condition is **predeterminedness**:

$$\mathbb{E}[x_{it}v_{is}] = 0 \quad \text{for } s \geq t,$$

which allows x_{it} to be correlated with past errors. Under the additional regularity conditions, this can guarantee consistency of the WG estimator:

$$\hat{\beta}_{WG} \xrightarrow{P} \beta$$

as $T \rightarrow \infty$.

However, for the FDOLS estimator, we still require the first differenced exogeneity assumption:

$$\mathbb{E}[\Delta x_{it} \Delta v_{it}] = 0$$

for consistency of the FDOLS estimator $\hat{\beta}_{FDOLS} \xrightarrow{P} \beta$ as $T \rightarrow \infty$.

Because the large- T consistency result for WG can use weaker sequential-exogeneity conditions, this comparison favours WG on robustness grounds; efficiency and the serial-correlation structure may still favour FDOLS in particular

applications.

Proposition 613 (WG and FDOLS Estimators in Large T)

Under suitable weak-dependence, moment, and identification conditions, the WG estimator is consistent under predeterminedness in large- T samples.

The FDOLS still requires the first-differenced exogeneity assumption.

13.4 Random Effects Model

We now go back to the scenario of small T panels whilst assuming random effects and strictly exogenous.

Assumption 614 (Strict Exogeneity)

All K explanatory variables are uncorrelated with individual time-varying error terms over all periods:

$$\mathbb{E}[\mathbf{x}_{it}v_{is}] = \mathbf{0} \quad (537)$$

for all $s, t \in \{1, 2, \dots, T\}$.

Assumption 615 (Uncorrelated Individual Effects/Random Effects)

All K explanatory variables $\mathbf{x}_{it}$ is uncorrelated with all individual-specific effects η_i :

$$\mathbb{E}[\mathbf{x}_{it}\eta_i] = \mathbf{0} \quad (538)$$

for $t = 1, 2, \dots, T$.

Notice that this is a very nice setting to be in! We have the strongest form of exogeneity and we have random effects. So what is the issue then? The issue is that the previous estimators we have seen, the within-groups and the first differenced OLS is **not** efficient if we have random effects:

$$\mathbb{E}[\mathbf{x}_{it}\eta_i] = \mathbf{0}.$$

Recall that these two estimators were constructed under the assumption of fixed effects $\mathbb{E}[\mathbf{x}_{it}\eta_i] \neq \mathbf{0}$ and as a result, we had to perform transformation operations on the levels of the series $(y_{it}, \mathbf{x}_{it})$ to get rid of η_i . This then allowed us in those scenarios to estimate β under fixed effects. However, if you think about it, if we are under the assumption of random effects $\mathbb{E}[\mathbf{x}_{it}\eta_i] = \mathbf{0}$, then there is no need to perform such a transformation to get rid of η_i and in fact, we are losing information as we could've just used the levels. We do not have **efficient estimators**. Therefore, we figure out how to construct efficient estimators under this assumption of random effects.

We could also use pooled OLS since contemporaneous exogeneity $\mathbb{E}[\mathbf{x}_{it}v_{it}] = \mathbf{0}$ and random effects $\mathbb{E}[\mathbf{x}_{it}\eta_i] = \mathbf{0}$ both hold (strict exogeneity implies contemporaneous exogeneity). However, pooled OLS is not efficient under the standard error-components covariance structure because the errors $u_{it} = \eta_i + v_{it}$ are serially correlated through the time-invariant error η_i .

To this end, we will now look at another estimator called the between groups estimator. However, we will later see that this estimator is not efficient either but once we established this estimator, we will then be able to find an efficient estimator.

13.4.1 Between Groups Estimator

We continue assuming strict exogeneity and random effects. We now define the between groups transformation.

Definition 616 (Between Transformation)

The between transformation is given by taking the mean of observations:

$$\bar{y}_i = \frac{1}{T} \sum_{t=1}^T y_{it} \quad (539)$$

$$\bar{x}_i = \frac{1}{T} \sum_{t=1}^T x_{it} \quad (540)$$

$$\bar{v}_i = \frac{1}{T} \sum_{t=1}^T v_{it} \quad (541)$$

$$\bar{\eta}_i = \frac{1}{T} \sum_{t=1}^T \eta_{it} = \eta_i. \quad (542)$$

The between transformation structural equation is:

$$\bar{y}_i = \bar{x}_i \beta + (\bar{v}_i + \eta_i) \quad (543)$$

for $i = 1, 2, \dots, N$.

Remark 617. Again notice that the time-invariant error term has the mean:

$$\bar{\eta}_i = \frac{1}{T} \sum_{t=1}^T \eta_{it} = \frac{1}{T} T \eta_i = \eta_i.$$

In contrast, the individual specific time-varying error term has the mean:

$$\bar{v}_i = \frac{1}{T} \sum_{t=1}^T v_{it} = \frac{1}{T} (v_{i1} + v_{i2} + \dots + v_{iT}).$$

Notice the similarity to the within groups transformation which demeaned the variables:

$$\tilde{y}_{it} = y_{it} - \bar{y}_i = y_{it} - \frac{1}{T} \sum_{s=1}^T y_{is}.$$

These transformations are in fact orthogonal complements to one another. We can now define the between groups estimator.

Definition 618 (Between Groups Estimator)

Suppose that the between groups structural equation is:

$$\bar{y}_i = \bar{x}_i \beta + (\bar{v}_i + \eta_i) \quad (544)$$

for $i = 1, 2, \dots, N$.

The **between groups estimator** is the OLS estimator of the between groups structural equation:

$$\hat{\beta}_{BG} = (\bar{X}^T \bar{X})^{-1} (\bar{X})^T \bar{y}. \quad (545)$$

Notice that in our structural equation:

$$\bar{y}_i = \bar{x}_i \beta + (\bar{v}_i + \eta_i)$$

we have $\bar{\mathbf{x}}_i = \frac{1}{T} [\mathbf{x}_{i1} + \dots + \mathbf{x}_{iT}]$ and $\bar{v}_i = \frac{1}{T} [v_{i1} + \dots + v_{iT}]$. Therefore, in order for our usual orthogonality assumption in OLS between the regressor and error term, we require strict exogeneity $\mathbb{E}[\mathbf{x}_{it}v_{is}] = \mathbf{0}$ for all $s, t \in \{1, \dots, T\}$. Furthermore, notice that we did not get rid of the time-invariant variable η_i and therefore we need to assume random effects so that $\mathbb{E}[\mathbf{x}_{it}\eta_i] = \mathbf{0}$. With these assumptions, we have consistency of this estimator.

Proposition 619 (Consistency of Between-Groups OLS)

Suppose that we have strict exogeneity and random effects. Then the between-groups OLS estimator is a consistent estimator of $\boldsymbol{\beta}$:

$$\hat{\boldsymbol{\beta}}_{BG} \xrightarrow{P} \boldsymbol{\beta} \quad (546)$$

as $N \rightarrow \infty$ with T fixed.

However, an issue is that even under these assumptions and this transformation, $\hat{\boldsymbol{\beta}}_{BG}$ is **not efficient**. The reason for this is that we are losing information when we average the observations over the T time periods which discards information in the time series variation. However, we know that these time series variations can be useful for estimating $\boldsymbol{\beta}$ since this is what we used as the source of our transformations when we constructed within groups and first-differenced OLS estimators.

So even though it seems like we are back at where we started, we shall soon see that we can use the between groups estimator in our process of constructing an efficient estimator when we have random effects and strict exogeneity.

13.4.2 Prereqs For Random Effects Generalised Least Squares Estimator

We have seen that the between-groups estimator is a consistent estimator but not efficient. The current issue is that the time-invariant error η_i is causing a lot of complications. Recall it was this term that caused serial correlation in the setting of pooled OLS. Also recall that it is because of this term that we applied the WG and FDOLS transformations to remove it; the BG transformation retains it but averages away within-individual variation. Therefore, in this section, we shall estimate the variance of this error term and apply what we saw in generalised least squares to deal with this inefficiency.

We will use both the between-groups estimator and the within-groups estimator to construct the random effect generalised least squares estimator. The big idea is that from the between groups estimator and the within groups estimator, we have the structural equations:

$$\bar{y}_i = \bar{\mathbf{x}}_i \boldsymbol{\beta} + (\bar{v}_i + \eta_i)$$

$$\tilde{y}_i = \tilde{\mathbf{x}}_i \boldsymbol{\beta} + \tilde{v}_i$$

where the first equation is the between groups structural equation and the second equation is the within groups structural equation. Now, we can see that in the between groups equation:

$$\bar{y}_i = \bar{\mathbf{x}}_i \boldsymbol{\beta} + (\bar{v}_i + \eta_i)$$

we can estimate the variance of the error term $\bar{v}_i + \eta_i$, which we denote by $\sigma_{\bar{u}_i}^2$ where we define the error term $\bar{u}_i \equiv \bar{v}_i + \eta_i$.

From the within groups equation:

$$\tilde{y}_i = \tilde{\mathbf{x}}_i \boldsymbol{\beta} + \tilde{v}_i$$

we can estimate the variance of $\tilde{v}_i$, which we denote by $\sigma_{\tilde{v}_i}^2$.

From this, if we have $\sigma_{\bar{u}_i}^2$ from the between groups equation and σ_v^2 from the within groups equation, we can then back out what the variance of η_i is, which we denote by σ_η^2 . We can then construct the random effects GLS estimator from these estimated variances. However, in order to back out the variance of η_i , we require an assumption about independence of the error terms.

Assumption 620 (I.I.D Error Terms)

The two errors components η_i and v_{it} are both i.i.d and uncorrelated with each other:

$$\mathbb{E}[\eta_i v_{it}] = 0 \quad (547)$$

for $t = 1, 2, \dots, T$.

Now, our goal for this section is to estimate the variance of η_i , which we denote by σ_η^2 . This will take a few steps so don't lose sight of the big picture.

Step One: Using Between Groups Estimator

The first step is to retrieve the residuals from between groups estimator:

$$\hat{\bar{u}}_i = (\widehat{\eta_i + \bar{v}_i}) = \bar{y}_i - \bar{x}_i \hat{\beta}_{BG} \quad (548)$$

where $\hat{\beta}_{BG}$ is the between groups estimator. We can then use these residuals $\hat{\bar{u}}_i$ to construct the estimator

$$\hat{\sigma}_{\bar{u}}^2 = \frac{\sum_{i=1}^N (\widehat{\eta_i + \bar{v}_i})^2}{N - K} = \frac{\hat{\bar{u}}^T \hat{\bar{u}}}{N - K} \quad (549)$$

$\hat{\sigma}_{\bar{u}}^2$ estimates the variance $\eta_i + \bar{v}_i$, whereby the variance has the form

$$\sigma_\eta^2 + \left(\frac{1}{T}\right)\sigma_v^2.$$

See slide 41 of lecture 2 Panel data for why it is this form. This means we can write:

$$\hat{\sigma}_{\bar{u}}^2 = \frac{\hat{\bar{u}}^T \hat{\bar{u}}}{N - K} = \hat{\sigma}_\eta^2 + \left(\frac{1}{T}\right)\hat{\sigma}_v^2$$

or just

$$\hat{\sigma}_{\bar{u}}^2 = \hat{\sigma}_\eta^2 + \left(\frac{1}{T}\right)\hat{\sigma}_v^2. \quad (550)$$

We will use this estimator $\hat{\sigma}_{\bar{u}}^2$ later on when constructing the random effects GLS estimator.

Step Two: Using Within Groups Estimator

We now want to construct an estimator of the variance of v_{it} , denoted by σ_v^2 . To do so, we need the residuals from within groups estimator:

$$\hat{\bar{v}}_{it} = \bar{y}_{it} - \bar{x}_{it} \hat{\beta}_{WG} \quad (551)$$

where $\hat{\beta}_{WG}$ is the within groups estimator. We have already discussed how to then construct an estimator of σ_v^2 using the residuals $\hat{\bar{v}}_{it}$ when discussing the within groups estimator in Eq. (514) which we recall:

$$\hat{\sigma}_v^2 = \frac{\sum_{i=1}^N \sum_{t=1}^T \hat{\bar{v}}_{it}^2}{N(T-1) - K} = \frac{\hat{\bar{v}}^T \hat{\bar{v}}}{N(T-1) - K}.$$

Step Three: Combining Two Variance Estimators

To recap, we first used the between groups estimator to construct the estimator $\hat{\sigma}_{\bar{u}}^2$ for $\sigma_{\bar{u}}^2$ where $\bar{u} = \eta_i + \bar{v}_i$. We then used the within groups estimator to construct the estimator $\hat{\sigma}_v^2$ for σ_v^2 . Recall that the goal of all this was to estimate σ_{η}^2 . We can now use these two variance estimators to compute the variance estimator for σ_{η}^2 by computing:

$$\hat{\sigma}_{\eta}^2 = \hat{\sigma}_{\bar{u}}^2 - \left(\frac{1}{T}\right)\hat{\sigma}_v^2. \quad (552)$$

We also know that this estimator is in fact a consistent estimator.

Proposition 621 (Consistency of Time-Invariant Error Variance)

The estimator $\hat{\sigma}_{\eta}^2$ is a consistent estimator of the variance σ_{η}^2 .

We have now finally achieved our goal of constructing the variance estimator $\hat{\sigma}_{\eta}^2$. We can now use this estimator in our ultimate goal of constructing an estimator of β that is efficient in the setting of strictly exogenous and random effects (uncorrelated individual effects). To do so, we shall use techniques that we saw in generalised least squares.

13.4.3 Random Effects Generalised Least Squares Estimator

Now, we shall discuss how to construct an efficient estimator in the setting of strictly exogenous variables and uncorrelated individual effects.

First, recall that the matrix $\mathbf{X}_i$ is a $T \times K$ matrix where each row is $\mathbf{x}_{it}$ (k explanatory variables for observation i at time t). Then, we can stack the $\mathbf{X}_i$ matrices vertically over individuals to get $\mathbf{X}$ which is a $NT \times K$ matrix.

Now, we state important assumptions that we will need (and have to memorise) for our upcoming estimator to be efficient.

Assumption 622 (Classical Assumptions)

The classical assumptions are:

1. $\eta_i | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT} \sim i.i.d(0, \sigma_{\eta}^2)$ over $i = 1, \dots, N$
2. $v_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT} \sim i.i.d(0, \sigma_v^2)$ over $i = 1, \dots, N$ and $t = 1, \dots, T$
3. η_i and v_{it} are independent conditional on $\mathbf{X}_i$
4. $\mathbb{E}[\eta_i | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}] = 0$
5. $\mathbb{E}[v_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}] = 0$ for each $t \in \{1, 2, \dots, T\}$

Now, if we look carefully at the last two assumptions in particular:

1. $\mathbb{E}[\eta_i | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}] = 0$
2. $\mathbb{E}[v_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}] = 0$ for each $t \in \{1, 2, \dots, T\}$

we can see that they are in fact stricter conditions than uncorrelated individuals effects and strictly exogenous covariates.

Proposition 623 (Strengthened Uncorrelated Individual Effects)

The classical assumption:

$$\mathbb{E}[\eta_i | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}] = 0 \quad (553)$$

implies uncorrelated individual effects (random effects):

$$\mathbb{E}[\mathbf{x}_{it}\eta_i] = 0. \quad (554)$$

Proof. Apply LIE:

$$\begin{aligned} \mathbb{E}[\mathbf{x}_{it}\eta_i] &= \mathbb{E}[\mathbb{E}[\mathbf{x}_{it}\eta_i | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}]] \\ &= \mathbb{E}[\mathbf{x}_{it}\mathbb{E}[\eta_i | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}]] \\ &= \mathbb{E}[\mathbf{x}_{it} \cdot 0] \\ &= 0. \end{aligned}$$

□

Proposition 624 (Strengthened Strictly Exogenous Covariates)

The classical assumption:

$$\mathbb{E}[v_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}] = 0 \quad (555)$$

for each $t \in \{1, 2, \dots, T\}$ implies strictly exogenous covariates

$$\mathbb{E}[\mathbf{x}_{it}v_{is}] = 0 \quad (556)$$

for each $t, s \in \{1, 2, \dots, T\}$.

Proof. Apply LIE:

$$\begin{aligned} \mathbb{E}[\mathbf{x}_{it}v_{is}] &= \mathbb{E}[\mathbb{E}[\mathbf{x}_{it}v_{is} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}]] \\ &= \mathbb{E}[\mathbf{x}_{it}\mathbb{E}[v_{is} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}]] \\ &= \mathbb{E}[\mathbf{x}_{it} \cdot 0] \\ &= 0. \end{aligned}$$

□

Now, under the classical assumptions, we look at the linear regression model where we stack the T time periods for each individual i to get:

$$\mathbf{y}_i = \mathbf{X}_i\boldsymbol{\beta} + \mathbf{u}_i$$

for $i = 1, \dots, N$ with $\mathbf{u}_i = \boldsymbol{\eta}_i\mathbf{j}_T + \mathbf{v}_i$. Here, we have that just like in the OLS setting:

$$\mathbb{E}[\mathbf{y}_i | \mathbf{X}_i] = \mathbf{X}_i\boldsymbol{\beta}$$

and by cross-sectional independence:

$$\mathbb{E}[\mathbf{y} | \mathbf{X}] = \mathbf{X}\boldsymbol{\beta}.$$

In this linear regression model, the efficient estimator of β is known as the **Random Effects Generalised Least Squares (GLS)** estimator, which takes into account for the serial correlation $u_{it} = \eta_i + v_{it}$ present which made our pooled OLS inefficient.

The final step before we need to do before constructing our random effects GLS estimator is to state under all these assumptions, what does the variance covariance matrix of the error terms $\mathbf{u}$, which is the sum of η_i and $\mathbf{v}_i$ errors, $\Omega = \mathbb{E}[\mathbf{u}\mathbf{u}^T | \mathbf{X}]$ look like.

Proposition 625 (Random Effects GLS Variance-Covariance Matrix of Errors)

The variance-covariance matrix of errors associated to the random effects GLS estimator is given by:

$$\Omega \equiv \mathbb{E}[\mathbf{u}\mathbf{u}^T | \mathbf{X}] = \begin{bmatrix} \Omega_i & 0 & \dots & 0 \\ 0 & \Omega_i & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & \Omega_i \end{bmatrix} \quad (557)$$

where Ω is a $NT \times NT$ matrix where each entry Ω_i is a $T \times T$ matrix.

Proof. There are a few steps to this proof. We will have to make repeated use of the classical assumptions. First, we can show that:

$$\begin{aligned} \mathbb{E}[u_{it} | \mathbf{X}_i] &= \mathbb{E}[\eta_i + v_{it} | \mathbf{X}_i] \\ &= \mathbb{E}[\eta_i | \mathbf{X}_i] + \mathbb{E}[v_{it} | \mathbf{X}_i] \\ &= 0 + 0 = 0. \end{aligned}$$

We can then compute the variance

$$\begin{aligned} \text{Var}[u_{it} | \mathbf{X}_i] &= \mathbb{E}[u_{it}^2 | \mathbf{X}_i] - \mathbb{E}[u_{it} | \mathbf{X}_i]^2 \\ &= \mathbb{E}[u_{it}^2 | \mathbf{X}_i] - 0 \\ &= \mathbb{E}[(\eta_i + v_{it})^2 | \mathbf{X}_i] \\ &= \mathbb{E}[\eta_i^2 + 2\eta_i v_{it} + v_{it}^2 | \mathbf{X}_i] \\ &= \mathbb{E}[\eta_i^2 | \mathbf{X}_i] + 2\mathbb{E}[\eta_i v_{it} | \mathbf{X}_i] + \mathbb{E}[v_{it}^2 | \mathbf{X}_i] \\ &= \sigma_\eta^2 + 0 + \sigma_v^2 \\ &= \sigma_\eta^2 + \sigma_v^2. \end{aligned}$$

We then compute the covariance:

$$\begin{aligned} \mathbb{E}[u_{it} u_{is} | \mathbf{X}_i] &= \mathbb{E}[(\eta_i + v_{it})(\eta_i + v_{is}) | \mathbf{X}_i] \\ &= \mathbb{E}[\eta_i^2 + \eta_i v_{is} + \eta_i v_{it} + v_{it} v_{is} | \mathbf{X}_i] \\ &= \mathbb{E}[\eta_i^2 | \mathbf{X}_i] + \mathbb{E}[\eta_i v_{is} | \mathbf{X}_i] + \mathbb{E}[\eta_i v_{it} | \mathbf{X}_i] + \mathbb{E}[v_{it} v_{is} | \mathbf{X}_i] \\ &= \sigma_\eta^2 + 0 + 0 + 0 \\ &= \sigma_\eta^2. \end{aligned}$$

Now, we can stack the error terms u_{it} into the vector $\mathbf{u}_i$ which is a $T \times 1$ vector:

$$\mathbf{u}_i = \eta_i \mathbf{j}_T + \mathbf{v}_i.$$

So then variance-covariance matrix of this error term $\mathbf{u}_i$ for an individual is:

$$\begin{aligned} \mathbb{E}[\mathbf{u}_i \mathbf{u}_i^T | \mathbf{X}_i] &= \begin{bmatrix} \text{Var}[u_{i1} | \mathbf{X}_i] & \text{Cov}[u_{i1}, u_{i2} | \mathbf{X}_i] & \dots & \text{Cov}[u_{i1}, u_{iT} | \mathbf{X}_i] \\ \text{Cov}[u_{i2}, u_{i1} | \mathbf{X}_i] & \text{Var}[u_{i2} | \mathbf{X}_i] & \dots & \text{Cov}[u_{i2}, u_{iT} | \mathbf{X}_i] \\ \dots & \dots & \dots & \dots \\ \text{Cov}[u_{iT}, u_{i1} | \mathbf{X}_i] & \text{Cov}[u_{iT}, u_{i2} | \mathbf{X}_i] & \dots & \text{Var}[u_{iT} | \mathbf{X}_i] \end{bmatrix} \\ &= \begin{bmatrix} \sigma_\eta^2 + \sigma_v^2 & \sigma_\eta^2 & \dots & \sigma_\eta^2 \\ \sigma_\eta^2 & \sigma_\eta^2 + \sigma_v^2 & \dots & \sigma_\eta^2 \\ \dots & \dots & \dots & \dots \\ \sigma_\eta^2 & \sigma_\eta^2 & \dots & \sigma_\eta^2 + \sigma_v^2 \end{bmatrix} \equiv \Omega_i \end{aligned}$$

where Ω_i is a $T \times T$ matrix.

Finally, we stack this across all individuals to get:

$$\mathbb{E}[\mathbf{u} \mathbf{u}^T | \mathbf{X}] = \begin{bmatrix} \Omega_i & 0 & \dots & 0 \\ 0 & \Omega_i & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & \Omega_i \end{bmatrix} \equiv \Omega. \quad (558)$$

Note that Ω is a $NT \times NT$ matrix where each entry Ω_i is a $T \times T$ matrix. □

We can now state the random effects GLS estimator.

Definition 626 (Random Effects Generalized Least Squares Estimator)

The Random Effects Generalized Least Squares estimator is:

$$\hat{\boldsymbol{\beta}}_{GLS} = (\mathbf{X}^T \Omega^{-1} \mathbf{X})^{-1} \mathbf{X}^T \Omega^{-1} \mathbf{y} \quad (559)$$

where Ω is a $NT \times NT$ matrix.

We finally get our efficient estimator!

Proposition 627 (Efficiency of Random Effects GLS Estimator)

Under the classical random-effects assumptions and with known Ω , the Random Effects Generalized Least Squares estimator is the efficient linear unbiased estimator.

We can also show that it is a consistent estimator.

Proposition 628 (Consistency of Random Effects GLS Estimator)

Under strict exogeneity, uncorrelated individual effects, cross-sectional independence, and the usual moment and full-rank conditions, $\hat{\boldsymbol{\beta}}_{GLS}$ is a consistent estimator as $N \rightarrow \infty$ with T fixed.

The issue we have now is that Ω is a $NT \times NT$ matrix, which is really large. Trying to invert this matrix is computationally intensive. Therefore, we need to use other techniques instead to compute this random effects GLS estimator. We move onto another transformation that we can use that uses an OLS but gives the same result as the random effects GLS estimator.

13.4.4 Theta Differencing

We have argued that it is computationally intensive to compute the random effects generalised least squares. We can instead use theta-differencing transformation and use the standard OLS in this transformed model. First, we define what is this new transformation.

Definition 629 (Theta Differencing Transformation)

The theta differencing transformation is:

$$y_{it}^* = y_{it} - (1 - \theta)\bar{y}_i \quad (560)$$

where θ is a parameter that is the ratio of variances such that:

$$\theta^2 = \frac{\sigma_v^2}{\sigma_v^2 + T\sigma_\eta^2} \quad (561)$$

$$\bar{y}_i = \frac{1}{T} \sum_{s=1}^T y_{is} \quad (562)$$

where σ_v^2 is the variance of v_{it} error and σ_η^2 is the variance of η_i error.

If we know σ_v^2 and σ_η^2 , we can first compute θ^2 , then θ and $(1 - \theta)$. We can then retrieve the transformed variables to get our theta-differenced structural equation.

We can use this theta differencing to construct our new transformed structural equation.

Definition 630 (Theta Differencing)

Suppose that the theta differenced structural equation is:

$$y_{it}^* = \mathbf{x}_{it}^* \boldsymbol{\beta} + u_{it}^* \quad (563)$$

for $i = 1, 2, \dots, N$ where $u_{it}^* = \eta_i^* + v_{it}^*$.

The theta-differencing estimator is the OLS estimator of the theta differenced structural equation:

$$\hat{\boldsymbol{\beta}}_\theta = ((\mathbf{X}^*)^T \mathbf{X}^*)^{-1} (\mathbf{X}^*)^T \mathbf{y}^*. \quad (564)$$

Notice here that the time-invariant error term

$$\eta_i^* = \eta_i - (1 - \theta)\eta_i = \theta\eta_i$$

does **not** eliminate the time-invariant component η_i from the error term of the transformed model. Therefore, we still require the uncorrelated individual effects assumption:

$$\mathbb{E}[\mathbf{x}_{it}\eta_i] = 0$$

for consistency of the OLS estimator since the error term η_i is still present. However, η_i is eliminated when $\theta = 0$.

Furthermore, we also have that the theta-differenced individual-specific error term is:

$$v_{it}^* = v_{it} - (1 - \theta)\bar{v}_i = v_{it} - (1 - \theta)\frac{1}{T}(v_{i1} + \dots + v_{iT}).$$

Since our theta differenced explanatory variable $\mathbf{x}_{it}^*$ also contains $\frac{1}{T} \sum_{s=1}^T \mathbf{x}_{is}$, we therefore require strict exogeneity between the explanatory variable and the individual specific error term:

$$\mathbb{E}[\mathbf{x}_{it} v_{is}] = 0$$

for $t, s \in \{1, \dots, T\}$.

Proposition 631 (Consistency of Theta-Differenced OLS Estimator)

Under strict exogeneity and the random-effects assumption, the OLS estimator after θ -differencing is a consistent estimator of $\boldsymbol{\beta}$:

$$\hat{\boldsymbol{\beta}}_{\theta} \xrightarrow{P} \boldsymbol{\beta} \quad (565)$$

as $N \rightarrow \infty$ with T fixed.

Don't forget that the purpose of this new estimator was that it is a more computationally efficient way to compute random effects GLS!

Proposition 632 (Equivalency of Theta-Differenced OLS and Random Effects GLS)

The random effects GLS estimator is equivalent to the theta-differenced OLS estimator:

$$\hat{\boldsymbol{\beta}}_{GLS} = \hat{\boldsymbol{\beta}}_{\theta}. \quad (566)$$

In practice, we don't actually know the variance of the error terms σ_{η}^2 and σ_v^2 , so therefore we need to use consistent estimators. This leads us to feasible GLS.

Definition 633 (Feasible Random Effects Generalised Least Squares)

Feasible random effects generalised least squares estimators is the same as random effects generalised least squares estimator but with consistent estimators of variance of the error terms σ_{η}^2 and σ_v^2 .

From the classical assumptions, we can obtain these consistent estimators of the error terms using within groups and between groups estimator. We can also show that feasible GLS is the same as GLS and therefore is also consistent.

Proposition 634 (Feasible GLS and GLS)

The feasible random effects GLS estimator is asymptotically equivalent to the random effects GLS estimator.

Therefore, under the assumptions of strict exogeneity and uncorrelated individual effects, the feasible random effects GLS estimator is a consistent estimator of $\boldsymbol{\beta}$.

We can show that the GLS estimator coincides with the OLS estimator and the WG estimator under certain settings.

Proposition 635 (GLS and Pooled OLS Without Individual Effects)

When $\sigma_{\eta}^2 = 0$, the GLS estimator and the pooled OLS estimator coincide:

$$\hat{\boldsymbol{\beta}}_{GLS} = \hat{\boldsymbol{\beta}}_{OLS}. \quad (567)$$

This has a nice interpretation because when $\sigma_\eta^2 = 0$, this says that there are no individual-specific error terms η_i . This then mean that there is no serial correlation in the error $u_{it} = v_{it} + \eta_i = v_{it}$. But recall that the serial correlation in u_{it} is what caused our pooled OLS estimator to be inefficient! Therefore, since we no longer have serial correlation in u_{it} , this means that our pooled OLS estimator will be efficient and in this case, coincide with the random effects GLS estimator!

We now look at the second case where we see that the random effects GLS estimator coincides with the within groups estimator.

Proposition 636 (GLS and WG Estimator and LSDV Estimator in Large T)

As $T \rightarrow \infty$, then the GLS estimator converges to the within groups estimator:

$$\hat{\beta}_{GLS} - \hat{\beta}_{WG} \xrightarrow{P} \mathbf{0}. \quad (568)$$

Furthermore, as the Least Squares Dummy Variables estimator is equivalent to the within groups estimator, this implies that the GLS estimator also converges to the Least Squares Dummy Variables estimator:

$$\hat{\beta}_{GLS} - \hat{\beta}_{LSDV} \xrightarrow{P} \mathbf{0}. \quad (569)$$

Finally, as the random effects GLS estimator is an efficient estimator, this implies that in large T samples, the within groups estimator β_{WG} and least squares dummy variables estimator β_{LSDV} are **efficient estimators** even under the assumption of uncorrelated individual effects.

Proof. First recall the theta differencing transformation, which also coincides with the random effects GLS estimator:

$$y_{it}^* = y_{it} - (1 - \theta)\bar{y}_i$$

where

$$\theta^2 = \frac{\sigma_v^2}{\sigma_v^2 + T\sigma_\eta^2}$$

$$\bar{y}_i = \frac{1}{T} \sum_{s=1}^T y_{is}$$

Here, as $T \rightarrow \infty$, we have that $\theta^2 \rightarrow 0$. This then implies that $\theta \rightarrow 0$. This then means our theta differenced transformation becomes:

$$y_{it}^* = y_{it} - (1 - \theta)\bar{y}_i = y_{it} - \bar{y}_i = \tilde{y}_{it}$$

where $\tilde{y}_{it}$ is the within groups transformation! Therefore our theta differencing transformation coincides with the within groups transformation!

□

Here, we see that as $T \rightarrow \infty$, the theta transformation approaches the within transformation because $1 - \theta \rightarrow 1$. This indicates that with large- T samples the efficiency difference between random-effects GLS and the within estimator vanishes under the maintained random-effects assumptions. The time average itself does not generally converge to zero; it is the GLS transformation weight that converges to the within-transformation weight. For small T , the within estimator is generally not efficient under uncorrelated individual effects, whereas correctly specified random-effects GLS is efficient. Since LSDV is algebraically equivalent to WG for the slope coefficients, the same large- T equivalence applies to LSDV.

13.5 Testing for Random Effects

So far, we have made use of the uncorrelated individual effects assumption multiple times:

$$\mathbb{E}[\mathbf{x}_{it}\eta_i] = 0.$$

We are now interested in testing whether if **some** of the K explanatory variables $\mathbf{x}_{it}$ are correlated with the individual specific error η_i .

We assume strict exogeneity for all the tests in this section.

Assumption 637 (Strict Exogeneity)

All K explanatory variables are uncorrelated with individual time-varying error terms over all periods:

$$\mathbb{E}[\mathbf{x}_{it}v_{is}] = \mathbf{0} \quad (570)$$

for all $s, t \in \{1, 2, \dots, T\}$.

Our test will be similar to what we did in instrumental variables. In particular, we shall compare the difference between estimators that are consistent under random effects against estimators that are not consistent under random effects. Therefore, if there is a big difference between these estimators, that means that there is correlated individual effects so random effects does not hold:

$$\mathbb{E}[\eta_i \mathbf{x}_i] \neq 0.$$

We now recall which estimators are consistent estimators under strict exogeneity and **correlated** individual effects (fixed effects):

1. Within Groups $\hat{\beta}_{WG}$
2. First Differenced OLS $\hat{\beta}_{FDOLS}$.

Remark 638. Recall that these estimators were consistent but not efficient under random effects.

We also recall which estimators are consistent estimators under strict exogeneity and **uncorrelated** individual effects (random effects):

1. Pooled OLS $\hat{\beta}_{OLS}$
2. Between Groups $\hat{\beta}_{BG}$
3. Random Effects GLS $\hat{\beta}_{GLS}$.

Therefore, we need to pick and compare an estimator from these two lists of estimators. We now state the hypothesis test that we will need.

Definition 639 (Hypothesis Test for Random Effects)

$$H_0 : \mathbb{E}[\eta_i | \mathbf{X}_i] = 0. \quad (571)$$

Remark 640. Notice that this is a stronger test compared to testing for uncorrelated individual effects $\mathbb{E}[\mathbf{X}_i \eta_i] = 0$.

We impose the classical assumptions previously mentioned.

Assumption 641 (Classical Assumptions)

The classical assumptions are:

1. $\eta_i | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT} \sim i.i.d(0, \sigma_\eta^2)$ over $i = 1, \dots, N$
2. $v_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT} \sim i.i.d(0, \sigma_v^2)$ over $i = 1, \dots, N$ and $t = 1, \dots, T$
3. η_i and v_{it} are independent conditional on $\mathbf{X}_i$
4. $\mathbb{E}[\eta_i | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}] = 0$
5. $\mathbb{E}[v_{it} | \mathbf{x}_{i1}, \mathbf{x}_{i2}, \dots, \mathbf{x}_{iT}] = 0$ for each $t \in \{1, 2, \dots, T\}$

We can now state the Hausman test.

Definition 642 (Hausman Test)

The vector of differences are defined to be:

$$\hat{\mathbf{q}} = \hat{\boldsymbol{\beta}}_{WG} - \hat{\boldsymbol{\beta}}_{GLS}. \quad (572)$$

The scalar Hausman test statistic is:

$$h = \hat{\mathbf{q}}^T \left[\text{Var}[\hat{\mathbf{q}}] \right]^{-1} \hat{\mathbf{q}} \xrightarrow{D} \chi^2(K) \quad (573)$$

under the null hypothesis $H_0 : \mathbb{E}[\eta_i | \mathbf{X}_i] = 0$.

Furthermore, the variance of the vector of differences is given by:

$$\text{Var}[\hat{\mathbf{q}}] = \text{Var}[\hat{\boldsymbol{\beta}}_{WG}] - \text{Var}[\hat{\boldsymbol{\beta}}_{GLS}]. \quad (574)$$

Remark 643. *The variance term makes use of the fact that GLS estimator is efficient and therefore will have a smaller variance.*

Something to note is that in instrumental variable testing, we required conditional homoskedasticity to use the Hausman test. However, we also require conditional homoskedasticity to compute the GLS estimator, so therefore this assumption is satisfied!

We can also construct another Hausman test where we instead use the between groups estimator.

Definition 644 (Alternative Hausman Test)

The vector of differences are defined to be:

$$\hat{\mathbf{q}}^* = \hat{\boldsymbol{\beta}}_{WG} - \hat{\boldsymbol{\beta}}_{BG}. \quad (575)$$

The scalar Hausman test statistic is:

$$h^* = \hat{\mathbf{q}}^{*T} \left[\text{Var}[\hat{\mathbf{q}}^*] \right]^{-1} \hat{\mathbf{q}}^* \xrightarrow{D} \chi^2(K) \quad (576)$$

under the null hypothesis $H_0 : \mathbb{E}[\eta_i | \mathbf{X}_i] = 0$.

Furthermore, the variance of the vector of differences is given by:

$$\text{Var}[\hat{\mathbf{q}}^*] = \text{Var}[\hat{\boldsymbol{\beta}}_{WG}] + \text{Var}[\hat{\boldsymbol{\beta}}_{BG}]. \quad (577)$$

Remark 645. Notice the variance term has a + rather than a -. This is due to the fact of the orthogonality between the within groups and between groups estimator.

Now, if we look at a single scalar parameter $\beta^k \in \boldsymbol{\beta}$, we therefore have that the vector of differences is now a scalar difference:

$$\hat{q}^k = \hat{\beta}_{WG}^k - \hat{\beta}_{GLS}^k.$$

The test statistic then simplifies to:

$$h^k = \frac{\left(\hat{\beta}_{WG}^k - \hat{\beta}_{GLS}^k\right)^2}{\text{Var}[\hat{\beta}_{WG}^k] - \text{Var}[\hat{\beta}_{GLS}^k]} \xrightarrow{D} \chi^2(1) \quad (578)$$

under the null hypothesis $H_0 : \mathbb{E}[\eta_i | \mathbf{X}_i] = 0$. Equivalently, define the signed statistic:

$$z^k = \frac{\hat{\beta}_{WG}^k - \hat{\beta}_{GLS}^k}{\sqrt{\text{Var}[\hat{\beta}_{WG}^k] - \text{Var}[\hat{\beta}_{GLS}^k]}} \xrightarrow{D} \mathcal{N}(0, 1), \quad (579)$$

under the null hypothesis $H_0 : \mathbb{E}[\eta_i | \mathbf{X}_i] = 0$. Its square satisfies $(z^k)^2 = h^k$ asymptotically.

13.6 Time-Invariant Explanatory Variables

We now want to include the time-invariant explanatory variable $\mathbf{w}_i$ back into our structural model. The within transformation removes $\mathbf{w}_i$ together with the individual effect, so its coefficient cannot be identified by the standard fixed-effects estimator without additional instruments or restrictions. The following is only a schematic first-stage setup; a complete estimator also requires explicit exclusion, exogeneity, and rank conditions.

The first-stage linear projection is:

$$\mathbf{w}_i = \bar{\mathbf{x}}_i \boldsymbol{\Pi} + \mathbf{z}_i \boldsymbol{\theta} + \mathbf{r}_i. \quad (580)$$

In the scalar single-instrument case, relevance requires:

$$\theta \neq 0. \quad (581)$$

With vector-valued variables, this is replaced by the corresponding full-rank condition on the excluded-instrument coefficient matrix.

Assumption 646 (Serially Uncorrelated Time-Varying Errors)

The time-varying error terms v_{it} are serially uncorrelated.

$$y_{it} = \mathbf{x}_{it}\boldsymbol{\beta} + \mathbf{w}_i\boldsymbol{\gamma} + (\eta_i + v_{it}) \quad (582)$$

for $i = 1, 2, \dots, N$ and $t = 1, 2, \dots, T$.

13.7 Models with Predetermined Covariates

Assumption 647 (Fixed Effects)

$$\mathbb{E}[\mathbf{x}_{it}\eta_i] \neq \mathbf{0}. \quad (583)$$

Definition 648 (Predetermined Explanatory Variable)

Suppose that the time-varying component is serially uncorrelated. Then, the explanatory variable $\mathbf{x}_{it}$ is **predetermined** with respect to the time-varying error component v_{it} if:

$$\mathbb{E}[\mathbf{x}_{it}v_{is}] = \mathbf{0} \text{ for } s \geq t \quad (584)$$

$$\mathbb{E}[\mathbf{x}_{it}v_{is}] \text{ may differ from } \mathbf{0} \text{ for } s \leq t - 1. \quad (585)$$

We could also write this as:

$$\mathbb{E}[\mathbf{x}_{it}v_{is}] = 0 \text{ for } s = t, t + 1, \dots, T \quad (586)$$

$$\mathbb{E}[\mathbf{x}_{it}v_{is}] \text{ may differ from } 0 \text{ for } s = 1, 2, \dots, t - 1. \quad (587)$$

This is a relaxation of the assumption of strict exogeneity $\mathbb{E}[\mathbf{x}_{it}v_{is}] = \mathbf{0}$ for $s, t \in \{1, \dots, T\}$.

13.7.1 First Differenced 2SLS Estimator

We need:

$$\mathbb{E}[\Delta\mathbf{x}_{it}\Delta v_{it}] = \mathbf{0}.$$

However, we have already argued that $\Delta\mathbf{x}_{it}$ is correlated with Δv_{it} . What we can do instead is to treat $\Delta\mathbf{x}_{it}$ as an endogenous variable like what we saw in instrumental variables. Then, we can try to find an instrument $\mathbf{z}_{it}$ for $\Delta\mathbf{x}_{it}$ such that:

1. The instrument $\mathbf{z}_{it}$ is **valid**: $\mathbb{E}[\mathbf{z}_{it}\Delta v_{it}] = \mathbf{0}$
2. The instrument $\mathbf{z}_{it}$ is **informative**: $\mathbf{z}_{it}$ is correlated with $\Delta\mathbf{x}_{it}$.

$$\Delta y_{it} = \Delta\mathbf{x}_{it}\boldsymbol{\beta} + \Delta v_{it} \quad (588)$$

for $i = 1, 2, \dots, N$ and $t = 3, \dots, T$.

$$\mathbb{E}[\mathbf{x}_{i,t-1}\Delta v_{it}] = \mathbf{0}.$$

This is an instrumental-variable orthogonality condition; it does not make $\Delta \mathbf{x}_{it}$ exogenous and therefore does not justify FDOLS. Stack the valid instruments in $\mathbf{Z}$ and let $\mathbf{P}_Z = \mathbf{Z}(\mathbf{Z}^T \mathbf{Z})^{-1} \mathbf{Z}^T$. The first-differenced 2SLS estimator is

$$\hat{\boldsymbol{\beta}}_{FD2SLS} = [(\Delta \mathbf{X})^T \mathbf{P}_Z (\Delta \mathbf{X})]^{-1} (\Delta \mathbf{X})^T \mathbf{P}_Z \Delta \mathbf{y}.$$

Under instrument validity, relevance, cross-sectional independence, and the usual moment and rank conditions,

$$\hat{\boldsymbol{\beta}}_{FD2SLS} \xrightarrow{P} \boldsymbol{\beta} \quad \text{as } N \rightarrow \infty \text{ with } T \text{ fixed.}$$