

Asymptotic Theory

Edicson Luna

University of Pennsylvania

Identification

Consistency

Asymptotic distribution

When do we not have Consistency?

Identification

The 3 steps of good work

1. **Step 1:** Define the effect of interest.
It implies asking “what if?”. It is contrary to just describing (“what is”).
2. **Step 2:** Identification of the target parameter.
Identification links the thought experiment and data.
3. **Step 3:** Statistical inference.
In practice, we only see a finite sample of the observables.
Here, we want to use asymptotic theory to talk about the real parameters.

Essential concept

In general,

- ▶ We have a question (target parameter)
- ▶ We have a model (assumptions)
- ▶ We have some data (we are empiricists)
- ▶ We want to use the model and data to answer the question

Identification is about making a *logically coherent* empirical conclusion.

Identification

In this class, our main objective will be identifying the parameter β in the linear regression $Y = X\beta$ (considering one or more explanatory variables).

But what is identification? In simple words, it is to be able to express the parameter in terms of the data after we assume a model (in this case, a linear model). But let's add a little formality and much more intuition.

Formal definition

Let P denote the true distribution of the observed data $(y_i, x_i)_{i=1}^n$. Denote by $\mathbf{P} = \{P_\theta : \theta \in \Theta\}$ a model for the distribution of the observed data. We assume $P \in \mathbf{P}$. In other words, we assume $\exists \theta \in \Theta$ such that $P_\theta = P$. We are interested in θ . (Θ is the set of all possible values of θ).

We define the set $\Theta_0(P) := \{\theta \in \Theta : P_\theta = P\}$ which is called “the identified set”. We say that θ is identified if $\Theta_0(P)$ is a singleton $\forall P \in \mathbf{P}$.

More on identification

Notice that identification has nothing to do with statistical inference.

sample $\rightarrow$ population $\rightarrow$ unobserved parameters

Identification is about the second arrow. Statistical inference is about using the sample to learn about the population (we will be back to this soon).

The second arrow is logically the first thing to consider: can we recover the population parameter when we know the population distribution?

Why asymptotics?

Notice that when we analyzed the expectation of our estimators, we used finite sample properties. Now we will work with large sample properties.

But, why are we working with asymptotic properties if we (almost) always have finite data? The reason is that obtaining unbiased estimators is usually not possible. That is why economists typically focus on what would happen if $n \rightarrow \infty$.

Economists agree that consistency (I will define it soon) is the minimal requirement for an estimator.

Consistency

Consistency: formal definition

Let $\hat{\theta}$ be an estimator of θ based on some data $(X_i)_{i=1}^n$. Then, $\hat{\theta}$ is a **consistent** estimator of θ if $\forall \epsilon > 0$,

$$\mathbb{P}(|\hat{\theta} - \theta| > \epsilon) \rightarrow 0 \quad \text{as } n \rightarrow \infty$$

We also say that $\hat{\theta}$ **converges in probability** to θ . If $\hat{\theta}$ is not consistent, we say that the estimator is **inconsistent**.

Notice that if we have a consistent estimator, $\hat{\theta} \xrightarrow{P} \theta$. This means that **having consistency implies that we can identify the parameter** (which is our final goal). That is why we want to talk about consistency.

Theorems

Remember that previously we derived $\beta = [\text{Var}(X)]^{-1}\text{Cov}(X, Y)$. The LLN and CLT we studied before are not completely useful for what comes now. The reason is the following.

- Consider a sequence $X_1, X_2, \dots, X_n$ of RVs that are i.i.d. Now, fix n and define a sequence of RVs defined as $W_2, \dots, W_n$ such that $W_i = (X_i - \bar{X}_n)^2$ where $\bar{X}_n$ is the average of the first n X 's.

Theorems

- ▶ For simplicity compare W_2 and W_3 . Notice that in both cases you have the term $\bar{X}_n$.
- ▶ Therefore, there is not complete independence between W_2 and W_3 .
- ▶ And so, since there is a degree of dependence in the sequence $W_1, W_2, \dots$ (they are not i.i.d.), we can not use our theorems. The good news is that there are some relaxed versions of the LLN and the CLT.

Relaxed theorems

Assume there is weak dependence in a sequence of RVs $W_1, W_2, \dots$

- ▶ **LLN (V2):** As long as $E[W_i^2] < \infty$, it is true that $\bar{W}_n \xrightarrow{p} E[W_i]$ as $n \rightarrow \infty$.
- ▶ **CLT (V2):** As long as $E[W_i^{2+\delta}] < \infty$ for $\delta > 0$, it is true that $\frac{\bar{W}_n - E[W_i]}{\text{se}(\bar{W}_n)} \xrightarrow{d} Z$ as $n \rightarrow \infty$. (Z is the standard normal distribution). This is the same as saying $\sqrt{n}(\bar{W}_n - E[W_i]) \xrightarrow{d} \mathcal{N}(0, V)$.

Slutsky for consistency (V1)

We also need a new theorem. There are two versions of Slutsky. So far, let's state the one we need now.

Slutsky's theorem. Let $(X_n)_{n=1}^N$ and $(Y_n)_{n=1}^N$ two sequences of RVs. Take c as a constant. If $X_n \xrightarrow{P} X$ and $Y_n \xrightarrow{P} c$, it is true that

- ▶ $X_n + Y_n \xrightarrow{P} X + c$
- ▶ $X_n Y_n \xrightarrow{P} Xc$
- ▶ If $c \neq 0$, $\frac{X_n}{Y_n} \xrightarrow{P} \frac{X}{c}$

In summary, Slutsky is about being able to multiply the expectations in the limit. We will use the second item.

CLM Assumptions

Let's quickly remember our CLM assumptions (in the model with more than one regressor)

1. Linearity in parameters. That is, the true relationship between y_i and x_i is

$$y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \dots + \beta_{k-1} x_{ik-1} + \epsilon_i$$

2. There is random sampling. That is, the observations are i.i.d.
3. The matrix $\mathbb{E}[X'X]$ has complete rank. (i.e. it's invertible).
4. Zero conditional mean, $\mathbb{E}[\epsilon|X] = 0$. This also may be said as $\mathbb{E}[\epsilon_i|x_i] = 0$

Consistency of the estimator

Theorem. Under assumptions 1 to 4, the OLS estimator is consistent.

Proof. Remember

$$\begin{aligned}\hat{\beta} &= (X'X)^{-1}X'Y = \beta + (X'X)^{-1}X'\epsilon \\ &= \beta + \left(\sum_{i=1}^n x_i x_i' \right)^{-1} \sum_{i=1}^n x_i \epsilon_i \\ &= \beta + \left(\frac{1}{n} \sum_{i=1}^n x_i x_i' \right)^{-1} \frac{1}{n} \sum_{i=1}^n x_i \epsilon_i\end{aligned}$$

Consistency of the estimator

By the LLN, assuming weak dependence

$$\frac{1}{n} \sum_{i=1}^n x_i x_i' \xrightarrow{p} \mathbb{E}[x_i x_i']$$

$$\frac{1}{n} \sum_{i=1}^n x_i e_i \xrightarrow{p} \mathbb{E}[x_i e_i]$$

Notice that $\mathbb{E}[x_i e_i] = 0$ by the LIE. Then (by Slutsky's theorem)

$$\left(\frac{1}{n} \sum_{i=1}^n x_i x_i' \right)^{-1} \frac{1}{n} \sum_{i=1}^n x_i e_i \xrightarrow{p} \mathbb{E}[x_i x_i']^{-1} \times 0 = 0$$

Relaxation of the assumption

Notice that $\mathbb{E}[\epsilon_i|x_i] = 0$ implies $\mathbb{E}[x_i\epsilon_i] = 0$, but not the contrary. That is, we could perfectly relax $\mathbb{E}[\epsilon_i|x_i] = 0$ and assume that $\mathbb{E}[x_i\epsilon_i] = 0$ (which is a weaker assumption since it is implied by the former).

Asymptotic distribution

Motivation

Notice that consistency is important since it tells us that our estimator goes in the right direction. We at least know that as we add observations (which should be i.i.d.) we are getting closer and closer to the real parameter (i.e. $\hat{\beta}$ is converging into β as n increases). As mentioned before, consistency is the minimum requirement for any estimator.

Nonetheless, notice that consistency does not allow us to perform statistical inference. That is why, now we focus on the sampling distribution of the OLS estimators.

Convergence in distribution

- **Convergence in distribution:** A sequence $(X_n)_{n=1}^{\infty}$ or RVs, with CDFs $(F_n)_{n=1}^{\infty}$. We say $X_n \xrightarrow{d} X$ if

$$\lim_{n \rightarrow \infty} F_n(x) = F(x)$$

Assumptions

We need to bring back assumption 5.

5. The error is constant given any value of X .

$$\text{Var}(\epsilon|X) = \sigma^2 I_n$$

Notice that previously we stated a 6th assumption about ϵ_i being normally distributed (and therefore y_i). Now, we will not need it. The reason as you may expect is that even if ϵ_i is not normally distributed, we will have normality by the CLT.

Slutsky for asymptotic distribution (V2)

Here, we will need another Slutsky theorem.

Slutsky's theorem. Let $(X_n)_{n=1}^{\infty}$ and $(Y_n)_{n=1}^{\infty}$ two sequences of RVs. Take c as a constant. If $X_n \xrightarrow{d} X$ and $Y_n \xrightarrow{p} c$, it is true that

- ▶ $X_n + Y_n \xrightarrow{d} X + c$
- ▶ $X_n Y_n \xrightarrow{d} Xc$
- ▶ If $c \neq 0$, $\frac{X_n}{Y_n} \xrightarrow{d} \frac{X}{c}$

Again, in the proof we followed the one in the middle.

Asymptotic normality of OLS

Theorem. Under assumptions 1 to 5,

$$\sqrt{n}(\hat{\beta} - \beta) \xrightarrow{d} \mathcal{N}(0, V)$$

where $V := \Sigma^{-1}\Omega\Sigma^{-1}$ with $\Sigma = \mathbb{E}[x_i x_i']$ and $\Omega = \mathbb{E}[X' \epsilon \epsilon' X]$

Asymptotic normality of OLS

Proof. Remember $\hat{\beta} = (\sum_{i=1}^n x_i x_i')^{-1} \sum_{i=1}^n x_i y_i$ which is the same as $(\frac{1}{n} \sum_{i=1}^n x_i x_i')^{-1} \frac{1}{n} \sum_{i=1}^n x_i y_i$. Replacing y_i is not difficult to see

$$\sqrt{n}(\hat{\beta} - \beta) = \left(\frac{1}{n} \sum_{i=1}^n x_i x_i' \right)^{-1} \frac{1}{\sqrt{n}} \sum_{i=1}^n x_i \epsilon_i$$

By the CLT

$$\frac{1}{\sqrt{n}} \sum_{i=1}^n x_i \epsilon_i \xrightarrow{d} \mathcal{N}(0, \text{Var}(X_i \epsilon_i) = \mathbb{E}[X' \epsilon \epsilon' X] = \Omega)$$

Asymptotic normality of OLS

Finally, by Slutsky (V2)

$$\left(\frac{1}{n} \sum_{i=1}^n x_i x_i'\right)^{-1} \frac{1}{\sqrt{n}} \sum_{i=1}^n x_i \epsilon_i \xrightarrow{d} \mathcal{N}(0, \Sigma^{-1} \Omega \Sigma^{-1}).$$

In conclusion,

$$\sqrt{n}(\hat{\beta} - \beta) \xrightarrow{d} \mathcal{N}(0, V)$$

Notice that in this case, we did not require the errors to have a normal distribution (as we did in hypothesis testing). The reason, as you already know, is the CLT which gives us the normality if $n \rightarrow \infty$. The Slutsky used here is another one which I will define in the next slide.

Natural estimator for variance

To carry out inference, the proposed consistent estimator for the asymptotic variance is of the form

$$\hat{V} := \hat{\Sigma}^{-1} \hat{\Omega} \hat{\Sigma}^{-1}$$

where

$$\hat{\Omega} := X' \epsilon \epsilon' X = \frac{1}{n} \sum_{i=1}^n e_i^2 x_i x_i'$$

$$\hat{\Sigma} := \frac{1}{n} \sum_{i=1}^n x_i x_i'$$

where e_i is the residual of the regression.

Natural estimator for variance

It is evidently that $\hat{\Sigma} \xrightarrow{P} \Sigma$ by the LLN. Then, as long as $\hat{\Omega} \xrightarrow{P} \Omega$, we use Slutsky (V1) and conclude that $\hat{V} \xrightarrow{P} V$.

$$\begin{aligned}\hat{\Omega} &= \frac{1}{n} \sum_{i=1}^n e_i^2 x_i x_i' = \frac{1}{n} \sum_{i=1}^n (y_i - x_i' \hat{\beta})^2 x_i x_i' \\ &= \frac{1}{n} \sum_{i=1}^n (\epsilon_i - x_i' (\hat{\beta} - \beta))^2 x_i x_i' \\ &= \frac{1}{n} \sum_{i=1}^n \epsilon_i^2 x_i x_i' - 2 \frac{1}{n} \sum_{i=1}^n (\hat{\beta} - \beta)' x_i \epsilon_i x_i' + \frac{1}{n} \sum_{i=1}^n [(\hat{\beta} - \beta)' x_i]^2 x_i x_i' \\ &\xrightarrow{P} \Omega\end{aligned}$$

Given that the second and third terms converge to zero (remember that $\hat{\beta}$ is a consistent estimator for β) and the first term converges to Ω by the LLN.

When do we not have Consistency?

Endogeneity

Recall that the proof of consistency of OLS relied on one key fact:

$$\frac{1}{n} \sum_{i=1}^n x_i \epsilon_i \xrightarrow{p} \mathbb{E}[x_i \epsilon_i] = 0$$

We say that x_i is **endogenous** if this fails, i.e.

$$\mathbb{E}[x_i \epsilon_i] = \text{Cov}(x_i, \epsilon_i) \neq 0.$$

If x_i is endogenous, going back to our proof,

$$\hat{\beta} = \beta + \left(\frac{1}{n} \sum_{i=1}^n x_i x_i' \right)^{-1} \frac{1}{n} \sum_{i=1}^n x_i \epsilon_i \xrightarrow{p} \beta + \Sigma^{-1} \mathbb{E}[x_i \epsilon_i] \neq \beta$$

So $\hat{\beta}$ is inconsistent: it does not converge to the true β .

Three classic sources of endogeneity

In applied work, endogeneity typically arises from one of three sources:

1. **Omitted variable bias.**
2. **Simultaneity** (reverse causality).
3. **Measurement error** in x_i .

We will go through each case, show algebraically why $\text{Cov}(x_i, \epsilon_i) \neq 0$, and give an example.

Case 1: Omitted variable bias

Suppose the true model is

$$y_i = \beta_0 + \beta_1 x_{i1} + \beta_2 x_{i2} + \epsilon_i, \quad \mathbb{E}[\epsilon_i | x_{i1}, x_{i2}] = 0$$

but x_{i2} is unobserved, so we estimate

$$y_i = \beta_0 + \beta_1 x_{i1} + u_i, \quad u_i := \beta_2 x_{i2} + \epsilon_i$$

Then

$$\text{Cov}(x_{i1}, u_i) = \beta_2 \text{Cov}(x_{i1}, x_{i2})$$

which is nonzero whenever $\beta_2 \neq 0$ and x_{i2} is correlated with x_{i1} . In that case

$$\text{plim } \hat{\beta}_1 = \beta_1 + \beta_2 \frac{\text{Cov}(x_{i1}, x_{i2})}{\text{Var}(x_{i1})}$$

Case 1: Example

Example. Consider the wage equation

$$\log(wage)_i = \beta_0 + \beta_1 educ_i + \beta_2 ability_i + \epsilon_i$$

$ability_i$ is typically unobserved by the econometrician.

If more able individuals also acquire more schooling, $Cov(educ_i, ability_i) > 0$. Since $\beta_2 > 0$ (ability raises wages), omitting ability biases $\hat{\beta}_1$ *upward*: part of the ability effect is wrongly attributed to education.

Case 2: Simultaneity

Suppose y_i and x_i are jointly determined:

$$y_i = \beta x_i + \epsilon_i, \quad x_i = \gamma y_i + \eta_i$$

Substituting the first equation into the second and solving for x_i ,

$$x_i = \gamma(\beta x_i + \epsilon_i) + \eta_i \implies x_i = \frac{\gamma \epsilon_i + \eta_i}{1 - \gamma \beta}$$

Then

$$\text{Cov}(x_i, \epsilon_i) = \frac{\gamma}{1 - \gamma \beta} \text{Var}(\epsilon_i)$$

which is nonzero whenever $\gamma \neq 0$, i.e. whenever y_i also feeds back into x_i .

Case 2: Example

Example. Supply and demand. Suppose we want to estimate a demand curve

$$q_i = \beta p_i + \epsilon_i$$

where ϵ_i is a demand shock. But price is set in equilibrium, so it also responds to ϵ_i through the supply side (e.g.

$$p_i = \gamma q_i + \eta_i).$$

Price and the demand shock move together, so $\text{Cov}(p_i, \epsilon_i) \neq 0$ and OLS cannot separately recover the demand curve from the supply curve.

Case 3: Measurement error

Suppose the true model uses the (unobserved) true regressor x_i^* :

$$y_i = \beta x_i^* + \epsilon_i, \quad \mathbb{E}[\epsilon_i | x_i^*] = 0$$

but we only observe $x_i = x_i^* + u_i$, with **classical measurement error**: u_i mean zero, and independent of x_i^* and ϵ_i .

Substituting $x_i^* = x_i - u_i$,

$$y_i = \beta(x_i - u_i) + \epsilon_i = \beta x_i + \underbrace{(\epsilon_i - \beta u_i)}_{\text{new error}}$$

Then,

$\text{Cov}(x_i, \epsilon_i - \beta u_i) = \text{Cov}(x_i^* + u_i, \epsilon_i - \beta u_i) = -\beta \text{Var}(u_i) \neq 0$,
if $\beta \neq 0$.

Case 3: Attenuation bias and example

Solving out the resulting bias,

$$\text{plim } \hat{\beta} = \beta \frac{\text{Var}(x_i^*)}{\text{Var}(x_i^*) + \text{Var}(u_i)}$$

Since $\frac{\text{Var}(x_i^*)}{\text{Var}(x_i^*) + \text{Var}(u_i)} \in (0, 1)$, $\hat{\beta}$ is biased toward zero. This is known as **attenuation bias**.

Example. Suppose we want to estimate the effect of cognitive ability on earnings, but we only observe a noisy test score as a proxy for true ability. The noise in the test score attenuates the estimated effect of ability toward zero.

Takeaway

In all three cases, the source of endogeneity is different, but the consequence is the same: $\text{Cov}(x_i, \epsilon_i) \neq 0$, so OLS is **inconsistent**.

This motivates the need for alternative estimation strategies (e.g. instrumental variables) that can restore consistency even when x_i is endogenous.