

Econometrics (Review)

Lecture 1 EE426 – 2/2014

Chayanee Chawanote

Finite sample properties of estimators

- a random sample: $\{Y_1, Y_2, \dots, Y_n\}$ drawn from a population distribution that depends on an unknown parameter θ .
- An estimator of θ is a rule that assigns each possible outcome of the sample a value of θ .
- E.g., a natural estimator of mean μ is an average of the random sample.
- More generally, an estimator W of a parameter θ can be expressed in a function h (rule): $W = h(Y_1, Y_2, \dots, Y_n)$
- When a particular set of numbers, $\{y_1, y_2, \dots, y_n\}$, is plugged into the function h , we obtain as estimate of θ : $w = h(y_1, y_2, \dots, y_n)$
- W is a random variable because it depends on the random sample: we obtain different random samples from the population, the value of W can change.

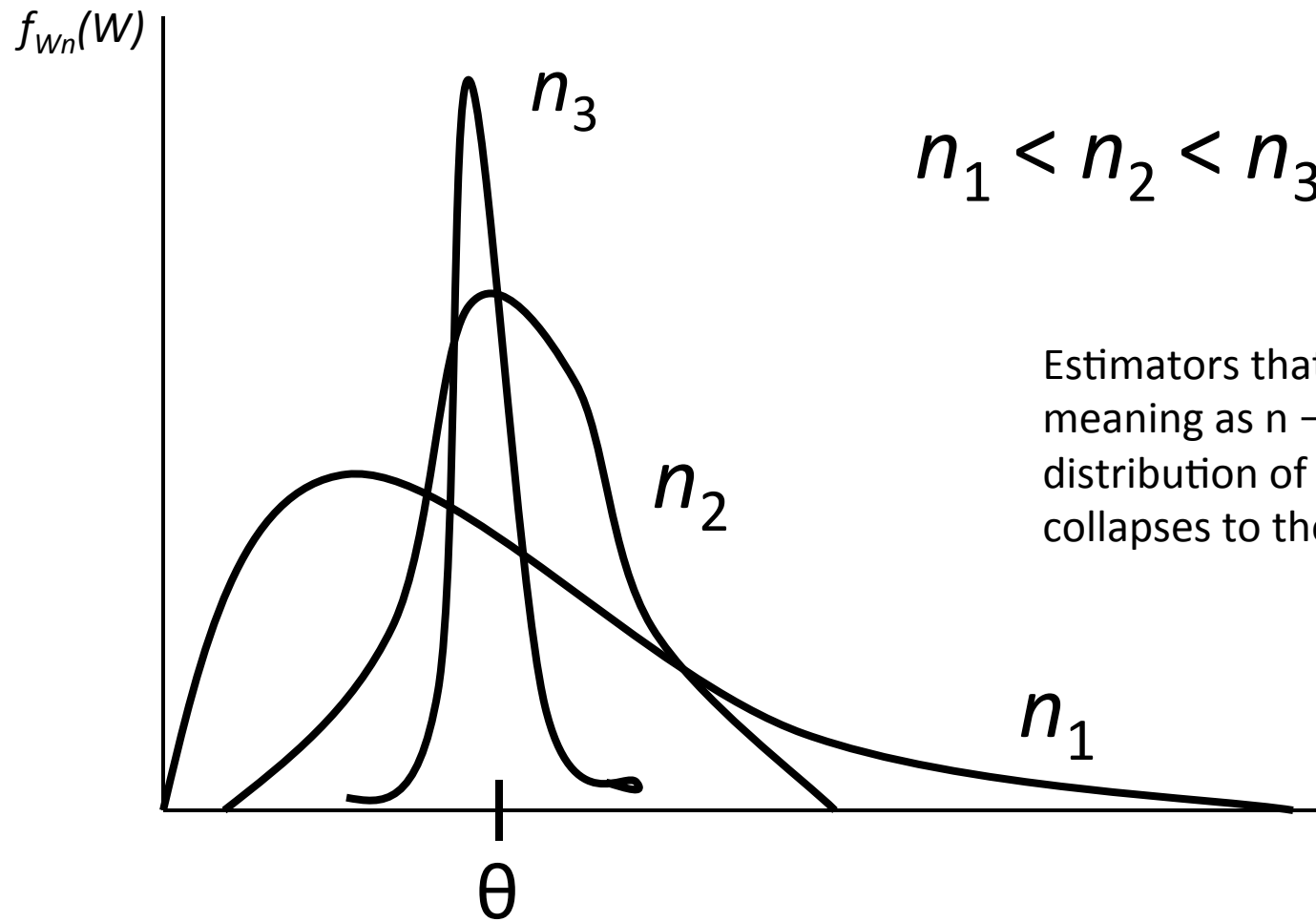
Finite sample properties of estimators

- **Unbiased estimator:** An estimator, W of θ , is an unbiased estimator if $E(W) = \theta$, for all possible values of θ
- Bias of an estimator: $\text{Bias}(W) \equiv E(W) - \theta$
- Sampling variance: the variance associated with a sampling distribution. It shows how spread out the distribution of an estimator is.
- If $\{Y_i; i = 1, 2, \dots, n\}$ is a random sample from a population with mean μ and variance σ^2 , then $\bar{Y}$ has the same mean as the population, but its sampling variance equal σ^2/n . Hence, the larger sample size, the smaller sampling variance.
- **Relative efficiency:** If W_1 and W_2 are two unbiased estimators, W_1 is efficient relative to W_2 when $\text{Var}(W_1) \leq \text{Var}(W_2)$ for all θ , with strict inequality for at least one value of θ

Asymptotic properties of estimators

- **Consistency**: Let W_n be an estimator of θ based on a sample $Y_1, Y_2, \dots, Y_n$ of size n . Then W_n is a consistent estimator of θ if for every $\varepsilon > 0$, $P(|W_n - \theta| > \varepsilon) \rightarrow 0$ as $n \rightarrow \infty$
- Probability limit: $\text{plim}(W_n) = \theta$
- If W_n is an unbiased estimator of θ and $\text{Var}(W_n) \rightarrow 0$ as $n \rightarrow \infty$, then $\text{plim}(W_n) = \theta$
- **Law of large number**: Let $Y_1, Y_2, \dots, Y_n$ be independent, identically distributed random variables with mean μ . Then, $\text{plim}(\bar{Y}_n) = \mu$
- **Asymptotic Normality** implies that $P(Z_n < z) \rightarrow \Phi(z)$ as $n \rightarrow \infty$, or $P(Z_n < z) \approx \Phi(z)$
 - The cumulative distribution function of Z_n gets closer and closer to the cdf of the standard normal distribution as the sample size n gets large.

Sampling Distributions as $n \uparrow$



$$n_1 < n_2 < n_3$$

Estimators that are consistent, meaning as $n \rightarrow \infty$, the distribution of the estimator collapses to the parameter value

Asymptotic properties of estimators

- **Central limit theorem**: the standardized average of any population with mean μ and variance σ^2 is asymptotically $\sim N(0,1)$, or

$$Z = \frac{\bar{Y} - \mu_Y}{\sigma / \sqrt{n}} \stackrel{a}{\sim} N(0,1)$$

- CLT basically says that for non-normal data, the distribution of the sample means has an approximate normal distribution, no matter what the distribution of the original data looks like, as long as the sample size is large enough (usually at least 30) and all samples have the same size.
- It's important because it means that we can approximate the distribution of certain statistics, even if we know very little about the underlying sampling distribution.

General approaches to parameter estimation

- We want estimators to be unbiasedness, consistency, and efficiency
- **Method of Moments**

The parameter θ is shown to be related to some expected value in the distribution of Y , usually $E(Y)$ or $E(Y^2)$.

Suppose θ is related to the population mean as $\theta = g(\mu)$ for some function g . Because the sample average $\bar{Y}$ is an unbiased and consistent estimator of μ , it gives us the estimator $g(\bar{Y})$ of μ . We replace the population moment, μ , with its sample counter part, $\bar{Y}$ when we know that $g(\bar{Y})$ is consistent for θ , and if $g(\mu)$ is a linear function of μ , then $g(\bar{Y})$ is unbiased as well.

- sample covariance
- sample correlation coefficient

General approaches to parameter estimation

- Maximum Likelihood

Let $\{Y_1, Y_2, \dots, Y_n\}$ be a random sample from the population distribution $f(y; \theta)$. Because of the random sampling assumption, the joint distribution of $\{Y_1, Y_2, \dots, Y_n\}$ is simply the product of the densities: $f(y_1; \theta)f(y_2; \theta) \cdots f(y_n; \theta)$. In the discrete case, this is $P(Y_1 = y_1, Y_2 = y_2, \dots, Y_n = y_n)$

Define the likelihood function as

$$L(\theta; Y_1, \dots, Y_n) = f(y_1; \theta)f(y_2; \theta) \cdots f(y_n; \theta)$$

The maximum likelihood estimator of θ , $\hat{\theta}$, is the value of θ that maximizes the likelihood function. This value depends on the random sample. Out of all the possible values for θ , the value that makes the likelihood of the observed data largest should be chosen.

General approaches to parameter estimation

- More convenient to work with the log-likelihood function:

$$\log[L(\theta; Y_1, \dots, Y_n)] = \sum_{i=1}^n \log[f(y_i; \theta)]$$

- Maximum likelihood estimation (MLE) is usually consistent and sometimes unbiased, generally the most asymptotically efficient estimator when the population model $f(y)$ is correctly specified.
- The MLE is sometimes the minimum variance unbiased estimator.
- If the distribution of Y conditional on a set of explanatory variables, $X_1, X_2, \dots, X_k$, then we replace the density with $f(Y_i | X_{i1}, \dots, X_{ik}; \theta_1, \dots, \theta_p)$, where this density is allowed to depend on p parameters.

General approaches to parameter estimation

- **Least Squares**

Find the m that makes the sum of squared deviations

$$\sum_{i=1}^n (Y_i - m)^2 \text{ as small as possible.}$$

- The sample mean, $\bar{Y}$, is a least squares estimator of the population mean.
- **For some important distributions**, including the normal and the Bernoulli, the sample average $\bar{Y}$ is also the maximum likelihood estimator of the population mean. All three approaches often result in the same estimator. In other cases, the estimators are similar but not identical.

Interval Estimation and Confidence Intervals

- Suppose the population has a normal($\mu, 1$) distribution and let $\{Y_1, Y_2, \dots, Y_n\}$ be a random sample from this population. Then, we have $\bar{Y} \sim \text{Normal}(\mu, 1/n)$.
- Standardize $\bar{Y}$, which has a standard normal distribution, we have

$$P(\bar{Y} - 1.96 / \sqrt{n} < \mu < \bar{Y} + 1.96 / \sqrt{n}) = 0.95$$

- The probability that the random interval $[\bar{Y} - 1.96 / \sqrt{n}, \bar{Y} + 1.96 / \sqrt{n}]$ contains the population mean μ is 95% >> interval estimate (w/ small y)
- It is also called a 95% confidence interval, meaning that this random interval (w/ big Y) contains μ with probability 0.95.
- This is NOT “the probability that μ is in the above interval is 0.95”.
- If unknown variance, we need the sample standard deviation which is constructed

$$s = \left(\frac{1}{n-1} \sum_{i=1}^n (y_i - \bar{y})^2 \right)^{1/2}$$

Interval Estimation and Confidence Intervals

- We must rely on the t distribution instead of the standard normal distribution.

$$\frac{\bar{Y} - \mu}{S / \sqrt{n}} \sim t_{n-1}$$

- So now we need c – the value such that 95% of the area in the t_{n-1} is between c and $-c$: $P(-c < t_{n-1} < c) = 0.95$. Once c has been properly chosen, we now have the random interval $[\bar{Y} - c \cdot S / \sqrt{n}, \bar{Y} + c \cdot S / \sqrt{n}]$ contains μ with prob .95.
- More generally, let c_α denote the $100(1-\alpha)$ percentile in the t_{n-1} distribution. Hence, a $100(1-\alpha)\%$ confidence interval is obtained as $[\bar{y} - c_{\alpha/2} \cdot s / \sqrt{n}, \bar{y} + c_{\alpha/2} \cdot s / \sqrt{n}]$
- The t distribution approaches the standard normal distribution as the df gets large: for $\alpha = .05$, $c_{\alpha/2} \rightarrow 1.96$ as $n \rightarrow \infty$.

Hypothesis Testing

- Mostly, we do a test about the mean in a normal population.
- **Null hypothesis** H_0 : $\mu = \mu_0$
- **Alternative hypothesis** H_a : $\mu > \mu_0$, $\mu < \mu_0$, or $\mu \neq \mu_0$
- Again, variance is replaced with the sample standard deviation. We need t-distribution to evaluate. Let call a **test statistic** as T . Under the null hypothesis, the random variable $T = \sqrt{n}(\bar{Y} - \mu_0) / S$ has a t_{n-1} distribution.
- Suppose we set 5% significance level. (Type I error: we reject H_0 when it is in fact true. We need fairly small Type I error.) The critical value c is chosen so that $P(T > c \mid H_0) = .05$ (one-tailed test). The rejection rule is $t > c$, where c is the 100(1- α) percentile in a t_{n-1} distribution.
- **P-value**: the largest significance level we could carry out the test and still fail to reject H_0

Matrix (basic)

EE426 Econometrics 2

Matrix Operation

• Matrix multiplication: $\mathbf{AB} = \left[\sum_{k=1}^n a_{ik} b_{kj} \right]$.

i^{th} row $\rightarrow$ $\begin{bmatrix} \mathbf{A} \\ a_{i1} & a_{i2} & a_{i3} & \dots & a_{in} \end{bmatrix} \begin{bmatrix} \mathbf{B} \\ b_{1j} \\ b_{2j} \\ b_{3j} \\ \vdots \\ b_{nj} \end{bmatrix} = \begin{bmatrix} \mathbf{AB} \\ \sum_{k=1}^n a_{ik} b_{kj} \end{bmatrix},$

$\uparrow$ j^{th} column $\uparrow$ $(i,j)^{\text{th}}$ element

- same result with summation between 2 variables

- Transpose: if $\mathbf{A} = [a_{ij}]_{m \times n}$, $\mathbf{A}' = [a_{ji}]_{n \times m}$
- Symmetric matrix: $\mathbf{A}' = \mathbf{A}$

Matrix Operation

- Properties of Matrix Operation

1. $(a + b)A = aA + bA$
2. $a(A + B) = aA + aB$
3. $(ab)A = a(bA)$
4. $a(AB) = (aA)B$
5. $A + B = B + A$
6. $(A + B) + C = A + (B + C)$
7. $(AB)C = A(BC)$
8. $A(B + C) = AB + AC$
9. $(A + B)C = AC + BC$
10. $IA = AI = A$
11. $A + 0 = 0 + A = A$
12. $A - A = 0$
13. $A0 = 0A = 0$
14. $AB \neq BA$

Note: AB is possible if #rows A = #columns B

- Properties of Transpose

1. $(A')' = A$
2. $(aA)' = aA'$
3. $(A + B)' = A' + B'$
4. $(AB)' = B'A'$
5. $x'x = \sum x^2$ when x is $n \times 1$ vector

Matrix Operation

- Partitioned Matrix Multiplication:

$A_{n \times k}$ and $B_{n \times m}$

$$\mathbf{A} = \begin{bmatrix} \mathbf{a}_1 \\ \mathbf{a}_2 \\ \vdots \\ \mathbf{a}_n \end{bmatrix}, \quad \mathbf{B} = \begin{bmatrix} \mathbf{b}_1 \\ \mathbf{b}_2 \\ \vdots \\ \mathbf{b}_n \end{bmatrix}$$

$$\mathbf{A}'\mathbf{B} = \sum_{i=1}^n \mathbf{a}_i' \mathbf{b}_i$$

Each i , $\mathbf{a}_i' \mathbf{b}_i$ is $k \times m$ matrix $\gg \mathbf{A}'\mathbf{B}$ can be written as the sum of n matrices

As a special case, we have

$$\mathbf{A}'\mathbf{A} = \sum_{i=1}^n \mathbf{a}_i' \mathbf{a}_i$$

Trace and Inverse

- Trace: for any $n \times n$ matrix A , $\text{tr}(A) = \sum_{i=1}^n a_{ii}$, is the sum of its diagonal elements.
- Properties of Trace
 1. $\text{tr}(I_n) = n$
 2. $\text{tr}(A') = \text{tr}(A)$
 3. $\text{tr}(A + B) = \text{tr}(A) + \text{tr}(B)$
 4. $\text{tr}(aA) = a\text{tr}(A)$
- Inverse: for any $n \times n$ matrix A has an inverse, A^{-1} , provided that $A^{-1}A = I_n$ and $AA^{-1} = I_n$. A is called “invertible or nonsingular matrix”
- Properties Inverse
 1. If an inverse exists, it is unique.
 2. $(aA)^{-1} = (1/a)A^{-1}$, $a \neq 0$ and A is invertible
 3. $(AB)^{-1} = B^{-1}A^{-1}$
 4. $(A')^{-1} = (A^{-1})'$

Linear Independence and Rank of Matrix

- Linear Independence: Let $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_r\}$ be a set of $n \times 1$ vector. These are linearly independent vectors if, and only if,

$$a_1\mathbf{x}_1 + a_2\mathbf{x}_2 + \dots + a_r\mathbf{x}_r = 0$$

implies that $a_1 = a_2 = \dots = a_r = 0$

- If the above has a not equal to 0, it means we can write at least one vector $\mathbf{x}$ as a linear combination of the others. Call $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_r\}$ as 'linearly dependent'.
 - $\text{Rank}(A)$ = maximum number of linearly independent columns of A .
 - If A ($n \times m$) has $\text{rank}(A) = m$, then A has full column rank
 - If A is $n \times m$, its rank can be at most m
- $\text{rank}(A') = \text{rank}(A)$
 - A ($n \times k$), $\text{rank}(A) \leq \min(n, k)$
 - If A ($k \times k$) and $\text{rank}(A) = k$, A is invertible.

Differentiation

- For a given $n \times 1$ vector $\mathbf{a}$, define the linear function as

$$f(\mathbf{x}) = \mathbf{a}'\mathbf{x}$$

- Partial derivatives: $\partial f(\mathbf{x})/\partial \mathbf{x} = \mathbf{a}'$

- A quadratic form in a symmetric matrix $\mathbf{A}$ ($n \times n$):

$$g(\mathbf{x}) = \mathbf{x}'\mathbf{A}\mathbf{x}$$

- Partial derivatives: $\partial g(\mathbf{x})/\partial \mathbf{x} = 2\mathbf{x}'\mathbf{A}$
($1 \times n$ vector)

Note: We can write a quadratic form in the summation as follows:

$$f(\mathbf{x}) = \mathbf{x}'\mathbf{A}\mathbf{x} = \sum_{i=1}^n a_{ii}x_i^2 + 2 \sum_{i=1}^n \sum_{j>1}^n a_{ij}x_i x_j$$

Random Vectors

- Expected value

1. Let y ($n \times 1$) be a random vector,

$$E(y) = [E(y_1), E(y_2), \dots, E(y_n)]'$$

2. Let Z ($n \times m$) be a random matrix,

$$E(Z) = [E(z_{ij})], \text{ (size } n \times m)$$

- Properties

1. $E(Ay + b) = AE(y) + b$, where A and b are nonrandom

2. $E(AZB) = AE(Z)B$, where $A(p \times n)$ and $B(m \times k)$ are nonrandom

Random Vectors

- Variance-Covariance Matrix

Let y ($n \times 1$) be a random vector. Then,

$$\text{Var}(y) = \begin{bmatrix} \sigma_1^2 & \sigma_{12} & \dots & \sigma_{1n} \\ \sigma_{21} & \sigma_2^2 & \dots & \sigma_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_{n1} & \sigma_{n2} & \dots & \sigma_n^2 \end{bmatrix} \quad \begin{aligned} \sigma_j^2 &= \text{Var}(y_j) \\ \sigma_{ij} &= \text{Cov}(y_i, y_j) \end{aligned}$$

is a symmetric matrix

1. $\text{Var}(a'y) = a'[\text{Var}(y)]a \geq 0$
2. $\text{Var}(y) = E[(y - \mu)(y - \mu)']$, $\mu = E(y)$
3. If y_i y_j are not correlated, $\text{Var}(y)$ is a diagonal matrix, $\text{Var}(y_i) = \sigma^2$ for $i = 1, 2, \dots, n$ and can be written as $\text{Var}(y) = \sigma^2 I_n$
4. $\text{Var}(Ay + b) = A[\text{Var}(y)]A'$