

- 到目前为止，我们讨论了多元线性回归模型的估计与推断：

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_k x_k + u$$

- 其中，我们主要关注**OLS**估计在给定任意样本情况下的性质，例如：

无偏性，方差，正态性等

- 这些对任何样本量（ $n > k+1$ ）都适用的性质叫做所谓有限样本性质、小样本性质或者精确性质。
- 这一章中，我们关心大样本（**large sample**）性质或者渐近（**asymptotic**）性质：样本量趋于无穷时**OLS**的性质

Chapter 5

多元回归分析：OLS的 渐近性

多元回归分析：OLS的渐近性

1. 线性
2. 随机抽样
iid,
ergodic
3. 无共线性
4. $E[u_i|x_i] = 0$
5. 同方差
正态

CLT版本 OLS有限样本性质

- 期望的无偏性：MLR.1 – MLR.4
- 方差公式：MLR.1 – MLR.5
- 高斯-马尔科夫定理：MLR.1 – MLR.5
- 正态性：MLR.1 – MLR.6

本章将讨论的OLS大样本性质

- 一致性：MLR.1 – MLR.4
- 渐近正态性：MLR.1 – MLR.5

(思考：为什么需要大样本性质?)

N 有限, $T \rightarrow \infty$ } 长面板, 向量
 N, T 时间序列 $N \geq 30, > 100$

$N \cdot T$

$N \rightarrow \infty, T$ 有限

短面板

$N \rightarrow \infty, T \rightarrow \infty$

① 无偏性很难维持, $E(\hat{\sigma}^2) = \sigma^2$ Jensen's

concave convex $E(\hat{\sigma}^2) \neq \sigma^2$ Jensen's 不等式,
 $E(f(x)) \neq f(E(x))$

② 大样本时, Central Limit Theorem

无需假设误差项的正态性!

$N \rightarrow \infty ?$

一致性

- 一致性 (**consistency**) 依概率收敛,

估计量 θ_n 是总体参数 θ 的一致估计, 如果随着 $n \rightarrow \infty$

$$P(|\theta_n - \theta| < \epsilon) \rightarrow 1 \text{ 对于任意 } \epsilon > 0.$$

另一种符号: $\text{plim } \theta_n = \theta$  估计量依概率收敛于真实总体值

- 解释:
 - 一致性意味着, 通过增加样本量, 可以使估计量任意接近 (arbitrarily close) 真实总体值的概率任意高 (arbitrarily high)。
 - 一个充分条件: $E(\theta_n)$ 等于或者趋于 θ , $\text{Var}(\theta_n)$ 趋于零
- 一致性是合理估计的最低要求

一致性

- **定理 5.1 (OLS的一致性)**

$$MLR.1-MLR.4 \Rightarrow \text{plim } \hat{\beta}_j = \beta_j, \quad j = 0, 1, \dots, k$$

- **简单回归模型的特例**

$$\text{plim } \hat{\beta}_1 = \beta_1 + \boxed{\text{Cov}(x_1, u)} / \text{Var}(x_1)$$

可以看出，如果解释变量是外生的，即与误差项不相关，则斜率估计是一致的。

- **假定 MLR.4'**

$$E(u) = 0$$

$$\text{Cov}(x_j, u) = 0$$

所有解释变量必须与误差项不相关。该假设弱于零条件平均假设MLR.4.

陈强

2. 随机序列的收敛

定义 随机序列 $\{x_n\}_{n=1}^{\infty} = \{x_1, x_2, x_3, \dots\}$ “依概率收敛”(converges in probability)于常数 a , 记为 $\text{plim}_{n \rightarrow \infty} x_n = a$, 或 $x_n \xrightarrow{p} a$, 如果 $\forall \varepsilon > 0$, 当 $n \rightarrow \infty$ 时, 都有 $\lim_{n \rightarrow \infty} P(|x_n - a| > \varepsilon) = 0$ 。

任意给定 $\varepsilon > 0$, 当 n 越来越大时, 随机变量 x_n 落在区间 $(a - \varepsilon, a + \varepsilon)$ 之外的概率收敛于 0。

$$x_n = \begin{cases} 0, & 1 - \frac{1}{n} \\ n, & \frac{1}{n} \end{cases}$$

$$E(x_n) = 0 \cdot (1 - \frac{1}{n}) + n \cdot \frac{1}{n} = 1$$

定义 随机序列 $\{x_n\}_{n=1}^{\infty}$ “依概率收敛”于随机变量 x ，记为 $x_n \xrightarrow{p} x$ ，如果随机序列 $\{x_n - x\}_{n=1}^{\infty}$ 依概率收敛于 0。

4

命题 (连续函数与依概率收敛可交换运算次序, preservation of convergence for continuous transformation) 假设 $g(\cdot)$ 为连续函数，

则 $\text{plim}_{n \rightarrow \infty} g(x_n) = g\left(\text{plim}_{n \rightarrow \infty} x_n\right)$ 。

$$\hat{\sigma}^2 \rightarrow \sigma^2$$

$$\sqrt{}$$

$$\sqrt{\sigma^2} \rightarrow \sigma$$

当 x_n 的分布越来越集中于某 x 附近时， $g(x_n)$ 的分布自然也就越来越集中于 $g(x)$ 附近。

$$\hat{\sigma}^2 \rightarrow \sigma^2, \quad \hat{\sigma} \rightarrow \sigma$$

概率收敛算子 $\text{plim}_{n \rightarrow \infty}$ 与连续函数 $g(\cdot)$ 可交换运算次序。期望算子 E 无此性质，一般 $E(x^2) \neq [E(x)]^2$ 。 Jensen.

3. 依均方收敛

定义 随机序列 $\{x_n\}_{n=1}^{\infty}$ “依均方收敛” (converges in mean square) 于常数 a , 如果 $\lim_{n \rightarrow \infty} E(x_n) = a$, $\lim_{n \rightarrow \infty} \text{Var}(x_n) = 0$ 。

命题 依均方收敛是依概率收敛的充分条件。

证明: 使用切比雪夫不等式(参见附录)。

当 x_n 的均值越来越趋于 a , 方差越来越小并趋于 0 时, 就有 $\text{plim}_{n \rightarrow \infty} x_n = a$, 即在极限处 x_n 退化(degenerate)为常数 a 。

此命题是依均方收敛概念的主要用途。

4. 依分布收敛

定义 记随机序列 $\{x_n\}_{n=1}^{\infty}$ 与随机变量 x 的累积分布函数(cdf)分别为 $F_n(\cdot)$ 与 $F(\cdot)$ 。

F_1 F_2

如果对于任意实数 c ，都有 $\lim_{n \rightarrow \infty} F_n(c) = F(c)$ ，则称随机序列 $\{x_n\}_{n=1}^{\infty}$ “依分布收敛” (converge in distribution) 于随机变量 x ，记为 $x_n \xrightarrow{d} x$ 。

x_n 离 x 很远

【例】当 t 分布的自由度越来越大时，其累积分布函数收敛于标准正态的累积分布函数。

依分布收敛

$$P(x=1) = P(x=0) = \frac{1}{2}, \quad \text{考虑 } x_n = 1 - x$$
$$P(x_n=1) = P(x_n=0) = \frac{1}{2}$$

如果 x 为正态分布，而 $x_n \xrightarrow{d} x$ ，则称 $\{x_n\}_{n=1}^{\infty}$ 为“渐近正态”
(asymptotically normal)。

依分布收敛意味着，两个随机变量的概率密度长得越来越像。

“依概率收敛”比“依分布收敛”更强(前者是后者的充分条件):

$$“x_n \xrightarrow{p} x” \Rightarrow “x_n \xrightarrow{d} x”$$

反之不然：当 x_n 与 x 的分布函数很接近时， x_n 与 x 的实际取值仍然可以很不相同(比如， x_n 与 x 相互独立)。

命题 假设 $g(\cdot)$ 为连续函数, 且 $x_n \xrightarrow{d} x$, 则 $g(x_n) \xrightarrow{d} g(x)$ 。

当 x_n 的分布越来越像 x 的分布时, $g(x_n)$ 的分布自然也越来越像 $g(x)$ 的分布。

例: 假设 $x_n \xrightarrow{d} z$, 其中 $z \sim N(0, 1)$,

则 $x_n^2 \xrightarrow{d} z^2$, 其中 $z^2 \sim \chi(1)$, 即 $x_n^2 \xrightarrow{d} \chi(1)$
(因为平方是连续函数)

正态分布

渐近标准正态的平方服从渐近 $\chi(1)$ 分布。

5.3 大数定律与中心极限定理

1. 弱大数定律(Weak Law of Large Numbers)

假定 $\{x_n\}_{n=1}^{\infty}$ 为独立同分布的随机序列, 且 $E(x_1) = \mu$, $\text{Var}(x_1) = \sigma^2$ 存在, 则样本均值 $\bar{x}_n \equiv \frac{1}{n} \sum_{i=1}^n x_i \xrightarrow{p} \mu$ 。

证明: 因为 $E(\bar{x}_n) = \mu$, 而

$$\text{Var}(\bar{x}_n) = \text{Var}\left(\frac{x_1 + \cdots + x_n}{n}\right) = \frac{1}{n^2} n \sigma^2 = \frac{\sigma^2}{n} \rightarrow 0, \text{ 故 } \bar{x}_n \text{ 依均方收敛于 } \mu.$$

因此, $\bar{x}_n \xrightarrow{p} \mu$ 。样本无限大时, 样本均值趋于总体均值, 故名“大数定律”。

$\bar{x}_n$ $\text{Var}(\bar{x}_n) = \frac{\sigma^2}{n}$ $\frac{\bar{x}_n - \mu}{\frac{\sigma}{\sqrt{n}}} \sim N(0, 1)$

2. 中心极限定理(Central Limit Theorem, 简记 CLT)

定理 假定 $\{x_n\}_{n=1}^{\infty}$ 为独立同分布的随机序列, 且 $E(x_1) = \mu$, $\text{Var}(x_1) = \sigma^2$ 存在, 则 $\sqrt{n}(\bar{x}_n - \mu) \xrightarrow{d} N(0, \sigma^2)$ 。

Normal

$\sqrt{n}(\bar{x}_n - \mu) \sim N(0, \sigma^2)$

根据弱大数定律, $(\bar{x}_n - \mu) \xrightarrow{p} 0$, 而 $\sqrt{n} \rightarrow \infty$, 故用 $\sqrt{n}(\bar{x}_n - \mu)$ (即 “ $\infty \cdot 0$ ” 型) 得到非退化分布。

ergodic stationary

进一步, $(\bar{x}_n - \mu)$ 收敛到 0 的速度与 $\frac{1}{\sqrt{n}}$ 收敛到 0 的速度类似 (二者乘积为非退化分布), 称为 “ $\sqrt{n}$ 收敛” (root-n convergence)。

收敛速度

直观上, 可视为 $\bar{x}_n \xrightarrow{d} N\left(\mu, \frac{\sigma^2}{n}\right)$; 但不严格, 因为 $\frac{\sigma^2}{n} \rightarrow 0$ 。

在一维情况下，中心极限定理可等价地写为 $\frac{\bar{x}_n - \mu}{\sqrt{\sigma^2/n}} \xrightarrow{d} N(0, 1)$,

但此形式不易推广到多维的情形。

推广到多维的情形：

假定 $\{\mathbf{x}_n\}_{n=1}^{\infty}$ 为独立同分布的随机向量序列，且 $E(\mathbf{x}_1) = \boldsymbol{\mu}$ ，
 $\text{Var}(\mathbf{x}_1) = \boldsymbol{\Sigma}$ 存在，则 $\sqrt{n}(\bar{\mathbf{x}}_n - \boldsymbol{\mu}) \xrightarrow{d} N(\mathbf{0}, \boldsymbol{\Sigma})$ 。

2. 一致估计量

定义 如果 $\text{plim}_{n \rightarrow \infty} \hat{\beta}_n = \beta$ ，则估计量 $\hat{\beta}_n$ 是参数 β 的一致估计量 (consistent estimator)。

一致性(consistency)意味着，当样本容量足够大时， $\hat{\beta}_n$ 依概率收敛到真实参数 β 。

这是对估计量最基本，也是最重要的要求。

如果估计方法不一致，意味着研究没有太大意义；因为无论样本容量多大，估计量也不会收敛到真实值。

3. 渐近正态分布与渐近方差

$$\hat{\beta} \xrightarrow{P} \beta$$

$$\text{Avar}(\hat{\beta}_n)$$
$$\sigma^2$$

定义 如果 $\sqrt{n}(\hat{\beta}_n - \beta) \xrightarrow{d} N(0, \Sigma)$, 其中 Σ 为半正定矩阵, 则称 $\hat{\beta}_n$ 为渐近正态分布(asymptotically normally distributed), 称 Σ 为渐近方差(asymptotic variance), 记为 $\text{Avar}(\hat{\beta}_n)$ 。 $\text{Avar}(\hat{\beta}_n) \cdot \sigma^2$

可近似地认为 $\hat{\beta}_n \xrightarrow{d} N(\beta, \Sigma/n)$ 。 $(\hat{\beta}_n - \beta)$ 收敛到 0 的速度与 $\frac{1}{\sqrt{n}}$ 收敛到 0 的速度相同, 称为“ $\sqrt{n}$ 收敛”(root-n convergence)。

4. 渐近有效

假设 $\hat{\beta}_n$ 与 $\tilde{\beta}_n$ 都是 β 的渐近正态估计量, 其渐近方差分别为 Σ 与 V 。如果 $(V - \Sigma)$ 为半正定矩阵, 则称 $\hat{\beta}_n$ 比 $\tilde{\beta}_n$ 更为渐近有效(asymptotically more efficient)。

5.5 渐近分布的推导

推导渐近分布的常用技巧，涉及依概率收敛与依分布收敛的交叉运算，统称“斯拉斯基定理” (Slutsky Theorem)。

$$(1) \quad \underbrace{x_n \xrightarrow{d} x}, \quad \underbrace{y_n \xrightarrow{p} a} \Rightarrow \underbrace{x_n + y_n \xrightarrow{d} x + a}。$$

在极限处， y_n 退化为常数 a ，故 $x_n + y_n$ 在极限处只是将 x_n 的渐近分布 x 位移到 $x + a$ 。

特例：如果 $a = 0$ ，则 $x_n + y_n \xrightarrow{d} x$ 。

$$(2) \quad \underbrace{x_n \xrightarrow{d} x}, \quad \underbrace{y_n \xrightarrow{p} 0} \Rightarrow \underbrace{x_n y_n \xrightarrow{p} 0}。$$

在极限处， y_n 退化为 0 ， x_n 有正常的渐近分布 x ，故 $x_n y_n$ 退化为 0 。

Ax $A \Sigma A'$

(3) 随机向量 $\mathbf{x}_n \xrightarrow{d} \mathbf{x}$, 随机矩阵 $\mathbf{A}_n \xrightarrow{p} \mathbf{A}$, $\mathbf{A}_n \mathbf{x}_n$ 可以相乘
 $\Rightarrow \mathbf{A}_n \mathbf{x}_n \xrightarrow{d} \mathbf{A} \mathbf{x}$ 。

特例：如果 $\mathbf{x} \sim N(\mathbf{0}, \Sigma)$, 则 $\mathbf{A}_n \mathbf{x}_n \xrightarrow{d} N(\mathbf{0}, \mathbf{A} \Sigma \mathbf{A}')$ 。

在极限处，随机矩阵 $\mathbf{A}_n$ 退化为常数矩阵 $\mathbf{A}$ 。

正态分布的线性组合仍服从正态分布，且
 $\text{Var}(\mathbf{A} \mathbf{x}) = \mathbf{A} \text{Var}(\mathbf{x}) \mathbf{A}' = \mathbf{A} \Sigma \mathbf{A}'$ 。

(4) 随机向量 $\mathbf{x}_n \xrightarrow{d} \mathbf{x}$, 随机矩阵 $\mathbf{A}_n \xrightarrow{p} \mathbf{A}$, $\mathbf{A}_n \mathbf{x}_n$ 可以相乘, $\mathbf{A}^{-1}$ 存在 $\Rightarrow$ 二次型 $\mathbf{x}_n' \mathbf{A}_n^{-1} \mathbf{x}_n \xrightarrow{d} \mathbf{x}' \mathbf{A}^{-1} \mathbf{x}$ 。

一致性

MLR1 $\rightarrow$ MLR4': consistency.

- 对于OLS的一致性, 只需要较弱的MLR.4'
- MLR.4'不成立则很有可能导致OLS的一致性不成立
- 遗漏变量的渐近偏误

MLR4: $E[u_i | \tilde{x}_i] = 0$
 $\downarrow$ $X_{i=1, \dots, N}$

MLR4': $E[u_i] = 0$
 $\text{Cov}(\tilde{x}_i, u_i) = 0$

$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + v$ $\leftarrow$ 真实模型

$E[xu] - E[x]E[u]$ 误设模型

$y = \beta_0 + \beta_1 x_1 + [\beta_2 x_2 + v] = \beta_0 + \beta_1 x_1 + u$

$E[E[xu|x]]$
 $= E[E[u|x]]$
 偏误
 $= 0$

$\Rightarrow \text{plim } \tilde{\beta}_1 = \beta_1 + \text{Cov}(x_1, u) / \text{Var}(x_1)$

$= \beta_1 + \beta_2 \text{Cov}(x_1, x_2) / \text{Var}(x_1) = \beta_1 + \beta_2 \delta_1$

MLR1 $\rightarrow$ MLR4'

如果遗漏变量与已含变量不相关, 则没有遗漏变量偏误

No CLT

渐近正态和大样本推断

渐近正态和大样本推断

- 在现实中，正态性假定MLR.6常常有问题
- 如果MLR.6不能成立，则t或F统计量可能是错的
- 幸运的是，即使没有MLR.6，OLS估计在大样本下也是正态分布的
- 这意味着，如果样本足够大，t或F检验仍然有效

MLR.5: 同方差 $Var(u_i | \bar{x}_i) = \sigma^2$

定理5.2 (OLS的渐近正态性 Asymptotic normality)

在假定MLR.1 – MLR.5下:

$$\frac{(\hat{\beta}_j - \beta_j)}{se(\hat{\beta}_j)} \underset{d}{\sim} \text{Normal}(0, 1)$$

在大样本中，标准化的估计量是正态分布的

$$\frac{\hat{u}'\hat{u}}{n-k} \xrightarrow{P} \sigma^2$$

而且 $\text{plim } \hat{\sigma}^2 = \sigma^2$

$$\hat{\sigma}^2: \frac{\hat{u}'\hat{u}}{n-k} = \frac{\sum \hat{u}_i^2}{n-k}$$

$$\frac{\hat{u}'\hat{u}}{n-k}$$

$$\sum \hat{u}_i^2$$

没有自相关, $x_i u_i, x_j u_j$, 无相关性. $i=1, \dots, N$

$g_i \equiv \bar{x}_i u_i$

$\bar{x}_i = \begin{bmatrix} x_{i1} \\ \vdots \\ x_{ik} \end{bmatrix}$

随机抽样

$j=1, \dots, N, i \neq j$

$\sqrt{n}(\bar{g}) \xrightarrow{d} N(0, E[\bar{g}_i \bar{g}_i'])$

CLT

$\hat{\beta} = \beta + (X'X)^{-1}(X'u)$

$X = \begin{bmatrix} x_{11} \dots x_{1k} \\ \vdots \\ x_{N1} \dots x_{Nk} \end{bmatrix}$
 $N \times K$

$= \beta + \left(\frac{1}{n} \sum x_i x_i' \right)^{-1} \left(\frac{1}{n} \sum x_i u_i \right)$

$X = \begin{bmatrix} x_1' \\ \vdots \\ x_n' \end{bmatrix}$

$X'X = \sum x_i x_i'$

$\xrightarrow{P} E(x_i x_i')^{-1} E(x_i u_i)$

MLR. 4'

$E[x_i u_i] = 0, E(u_i) = 0$
 $\text{Cov}(x_i, u_i) = 0$

$\frac{1}{\sqrt{n}} \sum x_i u_i \rightarrow N(0, E(\bar{g} \bar{g}'))$

CLT.

$$Ax \sim N(0, A \Sigma A').$$

$$\sqrt{n}(\hat{\beta} - \beta), \quad A \text{var}(\cdot)$$

$$g_i \equiv \tilde{x}_i u_i.$$

$$A \text{var}(\hat{\beta}) = E(x_i x_i')^{-1} E[g_i g_i'] E(x_i x_i')^{-1}$$

Robust Standard Error.

4.5.

同前?

$$E[\hat{u}_i^2 x_i x_i'] \quad \text{Sandwich,}$$

$$E[u_i^2 | x_i] = \sigma^2$$

$$E[E[u_i^2 x_i x_i' | x_i]]$$

$$\frac{1}{n} \sum \uparrow \rightarrow E[u_i x_i x_i']$$

$$= E[x_i x_i' \underbrace{E[u_i^2 | x_i]}_{\sigma^2}] = \sigma^2 E(x_i x_i')$$

$$\begin{aligned}
 \text{Avar}(\hat{\beta} - \beta) &= E(x_i x_i')^{-1} \cdot \underbrace{\sigma^2}_{\text{var}(u_i)} \cdot \underbrace{E(x_i x_i')}_{\text{cov}(x_i)}^{-1} \\
 &= \underline{\sigma^2 E(x_i x_i')^{-1}}
 \end{aligned}$$

$$\text{Var}(\hat{\beta} - \beta) = \sigma^2 (X'X)^{-1}$$

渐近正态和大样本推断

- 注意:

- MLR.1-MLR.5仍然是需要的, 例如, 异方差性

$$\sqrt{n}(\hat{\beta} - \beta)$$

- $\hat{\beta}_j - \beta_j$ 不是渐近正态! 根据一致性, $\hat{\beta}_j - \beta_j$ 与 $se(\hat{\beta}_j)$ 均趋于零。

- 可以证明 $\sqrt{n}(\hat{\beta}_j - \beta_j)$ 是渐近正态分布

- 与有限样本正态性对比:

大样本: 在假定MLR.1 – MLR.5下:

$$\frac{(\hat{\beta}_j - \beta_j)}{se(\hat{\beta}_j)} \underset{a}{\sim} \text{Normal}(0, 1)$$

有限样本: 在假定MLR.1 – MLR.6下:

$$\frac{\hat{\beta}_j - \beta_j}{se(\hat{\beta}_j)} \underset{n}{\sim} t_{n-k-1}$$

注意: 是否矛盾?

$$\hat{\beta} \xrightarrow{P} \beta, \quad \hat{\beta} \xrightarrow{d} \beta, \quad \sqrt{n}(\hat{\beta} - \beta) \xrightarrow{d} N$$

渐近正态和大样本推断

- 实际分析中

- 在大样本下，t分布接近于标准正态分布
- 即使没有MLR.6，t检验在大样本下也是有效的
- 置信区间与F检验也是有效的

$$F \rightarrow \chi^2(q)$$

q : 条件数.

- OLS抽样误差的渐近性分析

$$\widehat{Var}(\hat{\beta}_j) = \frac{\hat{\sigma}^2}{SST_j(1 - R_j^2)}$$

收敛到 $n \cdot Var(x_j)$ 收敛到 σ^2 收敛到一个固定值

渐近正态和大样本推断

- OLS抽样误差的渐近性分析（续）

$\widehat{Var}(\hat{\beta}_j)$ 以 $1/n$ 的速度收缩

$se(\hat{\beta}_j)$ 以 $\sqrt{1/n}$ 的速度收缩

- 这就是为什么大样本更好的原因

- 例子：出生体重方程中的标准误

$n = 1,388 \Rightarrow se(\hat{\beta}_{cigs}) = .00086$

$n = 694 \Rightarrow se(\hat{\beta}_{cigs}) = .0013$

只使用观测值的前半部分

$$\frac{\sqrt{\frac{1}{1388}}}{\sqrt{\frac{1}{694}}} = \frac{.00086}{.0013} \approx \sqrt{\frac{694}{1,388}}$$