

Econometrics I

(Tue., 8:50-10:20)

Room # 1 (法経講義棟)

- This class is based on **Statistics** (統計, 『コア・テキスト 統計学』 大屋 幸輔 著, 新世社) and **Econometrics** (計量経済, 『計量経済学』 山本 拓 著, 新世社), which are provided by Department of Economics, and **Basic Statistics** (統計基礎, 『コア・テキスト 統計学』 大屋 幸輔 著, 新世社), provided by Graduate School of Economics.

Thus, Statistics and Econometrics of undergraduate level are prerequisites.

- Furthermore, **Special Lectures in Economics (Statistical Analysis)** or **Statistical Analysis** (統計解析), provided by Graduate School of Economics, should be studied with this class.

Or, do self-study using the lecture notes of

<https://www2.econ.osaka-u.ac.jp/~tanizaki/class/2012/econome1/index.htm>

(The notes are written in English with Japanese translation for econometrics terms).

TA Session:

TAs: **Kaito Sakaguchi (坂口 開人)**
 u190954f [at] ecs. osaka-u. ac. jp

Jukina Hatakeyama (畠山 樹輝凪):
 u868710a [at] ecs. osaka-u. ac. jp

Date: **Thrs. 15:10-16:40**

Place: **Room #509 (509 セミナー室)**

Contents: **Basic Statistics, Matrix Algebra, and etc.**

Ask TAs directly in the TA session
 if you have questions about class, homework and etc.

- **Download the lecture notes from the following websites:**

<http://www2.econ.osaka-u.ac.jp/~tanizaki/class/2025/econome2/>

Contents

1	Maximum Likelihood Estimation (MLE, ^{さいゆう}最尤法)	9
2	Qualitative Dependent Variable (質的従属変数)	36
2.1	Discrete Choice Model (離散選択モデル)	37
2.1.1	Binary Choice Model (二値選択モデル)	37
2.2	Limited Dependent Variable Model (制限従属変数モデル)	58
2.3	Count Data Model (計数データモデル)	67
2.4	Applications	75
3	Panel Data	101
3.1	GLS — Review	101
3.2	Panel Model Basic	105

3.2.1	Fixed Effect Model (固定効果モデル)	106
3.2.2	Random Effect Model (ランダム効果モデル)	114
3.3	Hausman's Specification Error (特定化誤差) Test	118
3.4	Choice of Fixed Effect Model or Random Effect Model	120
3.4.1	The Case where X is Correlated with u — Review	120
3.4.2	Fixed Effect Model or Random Effect Model	122
3.5	Applications	124
4	Introduction to <u>Causal Inference</u> (因果推論)	
	— Micro-econometrics —	129
5	Generalized Method of Moments (GMM, 一般化積率法)	137
5.1	Method of Moments (MM, 積率法)	137
5.2	Generalized Method of Moments (GMM, 一般化積率法)	145

5.3	Generalized Method of Moments (GMM, 一般化積率法) II — Non-linear Case —	167
6	Time Series Analysis (時系列分析)	189
6.1	Introduction	189
6.2	Autoregressive Model (自己回帰モデル or AR モデル)	194
6.3	Unit Root (単位根) Test (Dickey-Fuller (DF) Test)	208
6.4	AR(p) model: Augmented Dickey-Fuller (ADF) Test	224
7	Bayesian Estimation	226
7.1	Elements of Bayesian Inference	226
7.1.1	Bayesian Point Estimate	228
7.1.2	Bayesian Interval for Parameter	230
7.1.3	Prior Probability Density Function	231

7.2	Sampling Methods	236
7.2.1	Gibbs Sampling	236
7.2.2	Metropolis-Hastings Algorithm	246

1 Maximum Likelihood Estimation (MLE, 最尤法^{さいゆうほう})

→ Review

1. The distribution function of $\{X_i\}_{i=1}^n$ is $f(x; \theta)$, where $x = (x_1, x_2, \dots, x_n)$.

θ is a vector or matrix of unknown parameters, e.g., $\theta = (\mu, \Sigma)$, where $\mu = E(X_i)$ and $\Sigma = V(X_i)$.

Note that X is a vector of random variables and x is a vector of their realizations (i.e., observed data).

Likelihood function $L(\cdot)$ is defined as $L(\theta; x) = f(x; \theta)$.

Note that $f(x; \theta) = \prod_{i=1}^n f(x_i; \theta)$ when $X_1, X_2, \dots, X_n$ are mutually independently and identically distributed.

The maximum likelihood estimate (MLE) of θ is the θ such that:

$$\max_{\theta} L(\theta; x). \quad \Longleftrightarrow \quad \max_{\theta} \log L(\theta; x).$$

Thus, MLE satisfies the following two conditions:

- (a) $\frac{\partial \log L(\theta; x)}{\partial \theta} = 0. \implies$ Solution of θ : $\tilde{\theta} = \tilde{\theta}(x)$
- (b) $\frac{\partial^2 \log L(\theta; x)}{\partial \theta \partial \theta'}$ is a negative definite matrix.

2. $x = (x_1, x_2, \dots, x_n)$ are used as the observations (i.e., observed data).

$X = (X_1, X_2, \dots, X_n)$ denote the random variables associated with the joint distribution $f(x; \theta) = \prod_{i=1}^n f(x_i; \theta)$.

3. Replacing x by X , we obtain the maximum likelihood **estimator** (MLE, which is the same word as the maximum likelihood **estimate**).

That is, MLE of θ satisfies the following two conditions:

- (a) $\frac{\partial \log L(\theta; X)}{\partial \theta} = 0. \implies$ Solution of θ : $\tilde{\theta} = \tilde{\theta}(X)$
- (b) $\frac{\partial^2 \log L(\theta; X)}{\partial \theta \partial \theta'}$ is a negative definite matrix.

4. **Fisher's information matrix** (フィッシャーの情報行列) or simply **information matrix**, denoted by $I(\theta)$, is given by:

$$I(\theta) = -E\left(\frac{\partial^2 \log L(\theta; X)}{\partial \theta \partial \theta'}\right),$$

where we have the following equality:

$$-E\left(\frac{\partial^2 \log L(\theta; X)}{\partial \theta \partial \theta'}\right) = E\left(\frac{\partial \log L(\theta; X)}{\partial \theta} \frac{\partial \log L(\theta; X)}{\partial \theta'}\right) = V\left(\frac{\partial \log L(\theta; X)}{\partial \theta}\right)$$

Note that $E(\cdot)$ and $V(\cdot)$ are expected with respect to X .

Proof of the above equality:

$$\int L(\theta; x) dx = 1$$

Take a derivative with respect to θ .

$$\int \frac{\partial L(\theta; x)}{\partial \theta} dx = 0$$

(We assume that (i) the domain of x does not depend on θ and (ii) the derivative $\frac{\partial L(\theta; x)}{\partial \theta}$ exists.)

(*) Differentiation of Composite Functions (合成関数の微分) or Chain rule (連鎖律):

$$\frac{\partial \log L(\theta; x)}{\partial \theta} = \frac{\partial \log L(\theta; x)}{\partial L(\theta; x)} \frac{\partial L(\theta; x)}{\partial \theta} = \frac{1}{L(\theta; x)} \frac{\partial L(\theta; x)}{\partial \theta}$$

i.e.,

$$\frac{\partial L(\theta; x)}{\partial \theta} = \frac{\partial \log L(\theta; x)}{\partial \theta} L(\theta; x)$$

Rewriting the above equation, we obtain:

$$\int \frac{\partial L(\theta; x)}{\partial \theta} dx = \int \frac{\partial \log L(\theta; x)}{\partial \theta} L(\theta; x) dx = 0,$$

i.e.,

$$E\left(\frac{\partial \log L(\theta; X)}{\partial \theta}\right) = 0.$$

Again, differentiating the above with respect to θ , we obtain:

$$\begin{aligned} & \int \frac{\partial^2 \log L(\theta; x)}{\partial \theta \partial \theta'} L(\theta; x) dx + \int \frac{\partial \log L(\theta; x)}{\partial \theta} \frac{\partial L(\theta; x)}{\partial \theta'} dx \\ &= \int \frac{\partial^2 \log L(\theta; x)}{\partial \theta \partial \theta'} L(\theta; x) dx + \int \frac{\partial \log L(\theta; x)}{\partial \theta} \frac{\partial \log L(\theta; x)}{\partial \theta'} L(\theta; x) dx \\ &= E\left(\frac{\partial^2 \log L(\theta; X)}{\partial \theta \partial \theta'}\right) + E\left(\frac{\partial \log L(\theta; X)}{\partial \theta} \frac{\partial \log L(\theta; X)}{\partial \theta'}\right) = 0. \end{aligned}$$

Therefore, we can derive the following equality:

$$-E\left(\frac{\partial^2 \log L(\theta; X)}{\partial \theta \partial \theta'}\right) = E\left(\frac{\partial \log L(\theta; X)}{\partial \theta} \frac{\partial \log L(\theta; X)}{\partial \theta'}\right) = V\left(\frac{\partial \log L(\theta; X)}{\partial \theta}\right),$$

where the second equality utilizes $E\left(\frac{\partial \log L(\theta; X)}{\partial \theta}\right) = 0$.

5. **Cramer-Rao inequality (クラメール・ラオの不等式)** is given by:

$$V(s(X)) \geq (I(\theta))^{-1},$$

where $s(X)$ denotes an unbiased estimator of θ .

$(I(\theta))^{-1}$ is called **Cramer-Rao Lower Bound (クラメール・ラオの下限)**.

Proof:

The expectation of $s(X)$ is:

$$E(s(X)) = \int s(x)L(\theta; x)dx.$$

Differentiating the above with respect to θ ,

$$\frac{\partial E(s(X))}{\partial \theta} = \int s(x) \frac{\partial L(\theta; x)}{\partial \theta} dx = \int s(x) \frac{\partial \log L(\theta; x)}{\partial \theta} L(\theta; x) dx$$

$$= \text{Cov}\left(s(X), \frac{\partial \log L(\theta; X)}{\partial \theta}\right)$$

For simplicity, let $s(X)$ and θ be scalars.

Then,

$$\begin{aligned} \left(\frac{\partial \mathbb{E}(s(X))}{\partial \theta}\right)^2 &= \left(\text{Cov}\left(s(X), \frac{\partial \log L(\theta; X)}{\partial \theta}\right)\right)^2 = \rho^2 \text{V}(s(X)) \text{V}\left(\frac{\partial \log L(\theta; X)}{\partial \theta}\right) \\ &\leq \text{V}(s(X)) \text{V}\left(\frac{\partial \log L(\theta; X)}{\partial \theta}\right), \end{aligned}$$

where ρ denotes the correlation coefficient between $s(X)$ and $\frac{\partial \log L(\theta; X)}{\partial \theta}$, i.e.,

$$\rho = \frac{\text{Cov}\left(s(X), \frac{\partial \log L(\theta; X)}{\partial \theta}\right)}{\sqrt{\text{V}(s(X))} \sqrt{\text{V}\left(\frac{\partial \log L(\theta; X)}{\partial \theta}\right)}}.$$

Note that $|\rho| \leq 1$.

Therefore, we have the following inequality:

$$\left(\frac{\partial E(s(X))}{\partial \theta}\right)^2 \leq V(s(X)) V\left(\frac{\partial \log L(\theta; X)}{\partial \theta}\right),$$

i.e.,

$$V(s(X)) \geq \frac{\left(\frac{\partial E(s(X))}{\partial \theta}\right)^2}{V\left(\frac{\partial \log L(\theta; X)}{\partial \theta}\right)}$$

Especially, when $E(s(X)) = \theta$, i.e., when $s(X)$ is an unbiased estimator of θ , the numerator of the right-hand side leads to one.

Therefore, we obtain:

$$V(s(X)) \geq \frac{1}{-E\left(\frac{\partial^2 \log L(\theta; X)}{\partial \theta^2}\right)} = (I(\theta))^{-1}.$$

Even in the case where $s(X)$ is a vector, the following inequality holds.

$$V(s(X)) \geq (I(\theta))^{-1},$$

where $I(\theta)$ is defined as:

$$\begin{aligned} I(\theta) &= -E\left(\frac{\partial^2 \log L(\theta; X)}{\partial \theta \partial \theta'}\right) \\ &= E\left(\frac{\partial \log L(\theta; X)}{\partial \theta} \frac{\partial \log L(\theta; X)}{\partial \theta'}\right) = V\left(\frac{\partial \log L(\theta; X)}{\partial \theta}\right). \end{aligned}$$

The variance of any unbiased estimator of θ is larger than or equal to $(I(\theta))^{-1}$.

Thus, $(I(\theta))^{-1}$ results in the lower bound of the variance of any unbiased estimator of θ .

6. Asymptotic Normality of MLE:

Let $\tilde{\theta}$ be MLE of θ .

As n goes to infinity, we have the following result:

$$\sqrt{n}(\tilde{\theta} - \theta) \longrightarrow N\left(0, \lim_{n \rightarrow \infty} \left(\frac{I(\theta)}{n}\right)^{-1}\right),$$

where it is assumed that $\lim_{n \rightarrow \infty} \left(\frac{I(\theta)}{n}\right)$ converges.

→ The proof will be shown later.

That is, when n is large, $\tilde{\theta}$ is approximately distributed as follows:

$$\tilde{\theta} \sim N\left(\theta, (I(\theta))^{-1}\right).$$

Suppose that $s(X) = \tilde{\theta}$.

When n is large, $V(s(X))$ is approximately equal to $(I(\theta))^{-1}$.

7. Optimization (最適化):

MLE of θ results in the following maximization problem:

$$\max_{\theta} \log L(\theta; x).$$

We often have the case where the solution of θ is not derived in closed form.

$\Rightarrow$ Optimization procedure

$$0 = \frac{\partial \log L(\theta; x)}{\partial \theta} \approx \frac{\partial \log L(\theta^*; x)}{\partial \theta} + \frac{\partial^2 \log L(\theta^*; x)}{\partial \theta \partial \theta'} (\theta - \theta^*).$$

Solving the above equation with respect to θ , we approximately obtain the following:

$$\theta = \theta^* - \left(\frac{\partial^2 \log L(\theta^*; x)}{\partial \theta \partial \theta'} \right)^{-1} \frac{\partial \log L(\theta^*; x)}{\partial \theta}.$$

Replace the variables as follows:

$$\theta \longrightarrow \theta^{(i+1)} \qquad \theta^* \longrightarrow \theta^{(i)}$$

Then, we have:

$$\theta^{(i+1)} = \theta^{(i)} - \left(\frac{\partial^2 \log L(\theta^{(i)}; x)}{\partial \theta \partial \theta'} \right)^{-1} \frac{\partial \log L(\theta^{(i)}; x)}{\partial \theta}.$$

⇒ **Newton-Raphson method (ニュートン・ラプソン法)**

Replacing $\frac{\partial^2 \log L(\theta^{(i)}; x)}{\partial \theta \partial \theta'}$ by $E \left(\frac{\partial^2 \log L(\theta^{(i)}; x)}{\partial \theta \partial \theta'} \right)$, we obtain the following optimization algorithm:

$$\begin{aligned} \theta^{(i+1)} &= \theta^{(i)} - \left(E \left(\frac{\partial^2 \log L(\theta^{(i)}; x)}{\partial \theta \partial \theta'} \right) \right)^{-1} \frac{\partial \log L(\theta^{(i)}; x)}{\partial \theta} \\ &= \theta^{(i)} + \left(I(\theta^{(i)}) \right)^{-1} \frac{\partial \log L(\theta^{(i)}; x)}{\partial \theta} \end{aligned}$$

⇒ **Method of Scoring (スコア法)**

Convergence speed might be improved, compared with Newton-Raphson method.

1. **Central Limit Theorem:** Let $X_1, X_2, \dots, X_n$ be mutually independently distributed random variables with mean $E(X_i) = \mu$ and variance $V(X_i) = \sigma^2 < \infty$ for $i = 1, 2, \dots, n$.

Define $\bar{X} = (1/n) \sum_{i=1}^n X_i$.

Then, the central limit theorem is given by:

$$\frac{\bar{X} - E(\bar{X})}{\sqrt{V(\bar{X})}} = \frac{\bar{X} - \mu}{\sigma / \sqrt{n}} \longrightarrow N(0, 1).$$

Note that $E(\bar{X}) = \mu$ and $V(\bar{X}) = \sigma^2/n$.

That is,

$$\sqrt{n}(\bar{X} - \mu) = \frac{1}{\sqrt{n}} \sum_{i=1}^n (X_i - \mu) \longrightarrow N(0, \sigma^2).$$

Note that $E(\bar{X}) = \mu$ and $nV(\bar{X}) = \sigma^2$.

In the case where X_i is a vector of random variable with mean μ and variance $\Sigma < \infty$, the central limit theorem is given by:

$$\sqrt{n}(\bar{X} - \mu) = \frac{1}{\sqrt{n}} \sum_{i=1}^n (X_i - \mu) \longrightarrow N(0, \Sigma).$$

Note that $E(\bar{X}) = \mu$ and $nV(\bar{X}) = \Sigma$.

2. **Central Limit Theorem II:** Let $X_1, X_2, \dots, X_n$ be mutually independently distributed random variables with mean $E(X_i) = \mu$ and variance $V(X_i) = \sigma_i^2$ for $i = 1, 2, \dots, n$.

Assume:

$$\sigma^2 = \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n \sigma_i^2 < \infty.$$

Define $\bar{X} = (1/n) \sum_{i=1}^n X_i$.

Then, the central limit theorem is given by:

$$\frac{\bar{X} - E(\bar{X})}{\sqrt{V(\bar{X})}} = \frac{\bar{X} - \mu}{\sigma / \sqrt{n}} \longrightarrow N(0, 1),$$

i.e.,

$$\sqrt{n}(\bar{X} - \mu) = \frac{1}{\sqrt{n}} \sum_{i=1}^n (X_i - \mu) \longrightarrow N(0, \sigma^2).$$

Note that $E(\bar{X}) = \mu$ and $nV(\bar{X}) \longrightarrow \sigma^2$.

In the case where X_i is a vector of random variable with mean μ and variance Σ_i , the central limit theorem is given by:

$$\sqrt{n}(\bar{X} - \mu) = \frac{1}{\sqrt{n}} \sum_{i=1}^n (X_i - \mu) \longrightarrow N(0, \Sigma),$$

where $\Sigma = \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n \Sigma_i < \infty$.

Note that $E(\bar{X}) = \mu$ and $nV(\bar{X}) \rightarrow \Sigma$.

[Review of Asymptotic Theories]

- **Convergence in Probability (確率収束)** $X_n \rightarrow a$, i.e., X converges in probability to a , where a is a fixed number.
- **Convergence in Distribution (分布収束)** $X_n \rightarrow X$, i.e., X converges in distribution to X . The distribution of X_n converges to the distribution of X as n goes to infinity.

Some Formulas

X_n and Y_n : Convergence in Probability

Z_n : Convergence in Distribution

- If $X_n \rightarrow a$, then $f(X_n) \rightarrow f(a)$.

- If $X_n \rightarrow a$ and $Y_n \rightarrow b$, then $f(X_n Y_n) \rightarrow f(ab)$.
- If $X_n \rightarrow a$ and $Z_n \rightarrow Z$, then $X_n Z_n \rightarrow aZ$, i.e., aZ is distributed with mean $E(aZ) = aE(Z)$ and variance $V(aZ) = a^2 V(Z)$.

[End of Review]

3. Weak Law of Large Numbers (大数の弱法則) — Review:

Suppose that $X_1, X_2, \dots, X_n$ are distributed.

As $n \rightarrow \infty$, $\bar{X} \rightarrow \lim_{n \rightarrow \infty} E(\bar{X})$ under $\lim_{n \rightarrow \infty} nV(\bar{X}) < \infty$, which is called the **weak law of large numbers**.

→ Convergence in probability

→ Proved by Chebyshev's inequality

- (i) Suppose that $X_1, X_2, \dots, X_n$ are assumed to be mutually independently and identically distributed with $E(X_i) = \mu$ and $V(X_i) = \sigma^2 < \infty$.

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

Then, $\bar{X} \longrightarrow \mu$ as $n \longrightarrow \infty$.

Note that $E(\bar{X}) = \mu$ and $nV(\bar{X}) = \sigma^2$.

- (ii) Suppose that $X_1, X_2, \dots, X_n$ are assumed to be mutually independently distributed with $E(X_i) = \mu_i$ and $V(X_i) = \sigma_i^2$.

Assume that

$$(a) \ E(\bar{X}) = \frac{1}{n} \sum_{i=1}^n \mu_i \longrightarrow \mu, \text{ i.e., } \lim_{n \rightarrow \infty} E(\bar{X}) = \mu, \text{ and}$$

$$(b) \ nV(\bar{X}) = \frac{1}{n} \sum_{i=1}^n \sigma_i^2 \longrightarrow \sigma^2 < \infty, \text{ i.e., } \lim_{n \rightarrow \infty} nV(\bar{X}) = \sigma^2 < \infty.$$

Then, $\bar{X} \rightarrow \mu$ as $n \rightarrow \infty$,

Note that $E(\bar{X}) = \frac{1}{n} \sum_{i=1}^n \mu_i$ and $nV(\bar{X}) = \frac{1}{n} \sum_{i=1}^n \sigma_i^2$.

- (iii) Suppose that $X_1, X_2, \dots, X_n$ are assumed to be serially correlated with $E(X_i) = \mu_i$ and $\text{Cov}(X_i, X_j) = \sigma_{ij}$.

Assume that

(a) $E(\bar{X}) = \frac{1}{n} \sum_{i=1}^n \mu_i \rightarrow \mu$, i.e., $\lim_{n \rightarrow \infty} E(\bar{X}) = \mu$, and

(b) $nV(\bar{X}) = \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^n \sigma_{ij} \rightarrow \sigma^2 < \infty$, i.e., $\lim_{n \rightarrow \infty} nV(\bar{X}) = \sigma^2 < \infty$.

Then, $\bar{X} \rightarrow \mu$ as $n \rightarrow \infty$,

Note that $E(\bar{X}) = \frac{1}{n} \sum_{i=1}^n \mu_i$ and $nV(\bar{X}) = \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^n \sigma_{ij}$.

4. Some Formulas of Expectation and Variance in Multivariate Cases

— Review:

A vector of random variable X : $E(X) = \mu$ and $V(X) \equiv E((X - \mu)(X - \mu)') = \Sigma$

Then, $E(AX) = A\mu$ and $V(AX) = A\Sigma A'$.

Proof:

$$E(AX) = AE(X) = A\mu$$

$$\begin{aligned} V(AX) &= E((AX - A\mu)(AX - A\mu)') = E(A(X - \mu)(A(X - \mu))') \\ &= E(A(X - \mu)(X - \mu)'A') = AE((X - \mu)(X - \mu)')A' = AV(X)A' = A\Sigma A' \end{aligned}$$

MLE: Asymptotic Properties

1. $X_1, X_2, \dots, X_n$ are random variables with density function $f(x; \theta)$.

Let $\hat{\theta}_n$ be a maximum likelihood estimator of θ .

Then, under some **regularity conditions**, $\hat{\theta}_n$ is a consistent estimator of θ and the asymptotic distribution of $\sqrt{n}(\hat{\theta} - \theta)$ is given by:

$$\sqrt{n}(\hat{\theta} - \theta) \longrightarrow N\left(0, \lim_{n \rightarrow \infty} \left(\frac{I(\theta)}{n}\right)^{-1}\right)$$

2. Regularity Conditions:

- (a) The domain of X_i does not depend on θ .
- (b) There exists at least third-order derivative of $f(x; \theta)$ with respect to θ , and their derivatives are finite.

3. Thus, MLE is

- (i) consistent,
- (ii) asymptotically normal, and
- (iii) asymptotically efficient.

Proof: The log-likelihood function is given by:

$$\log L(\theta) = \log \prod_{i=1}^n f(X_i; \theta) = \sum_{i=1}^n \log f(X_i; \theta)$$

X_i is a random variable.

Consider the distribution of

$$\frac{1}{n} \frac{\partial \log L(\theta)}{\partial \theta} = \frac{1}{n} \sum_{i=1}^n \frac{\partial \log f(X_i; \theta)}{\partial \theta}.$$

We have to obtain mean and variance of $\frac{\partial \log f(X_i; \theta)}{\partial \theta}$.

Suppose that X_i is a continuous type of random variable.

$f(x_i; \theta)$ denotes the density function.

Therefore, we have:

$$\int f(x_i; \theta) dx_i = 1$$

Taking the derivative with respect to θ on both sides, we obtain:

$$0 = \int \frac{\partial f(x_i; \theta)}{\partial \theta} dx_i = \int \frac{\partial \log f(x_i; \theta)}{\partial \theta} f(x_i; \theta) dx_i = E\left(\frac{\partial \log f(X_i; \theta)}{\partial \theta}\right)$$

Again, take the derivative with respect to θ on both sides as follows:

$$\begin{aligned} 0 &= \int \frac{\partial^2 \log f(x_i; \theta)}{\partial \theta \partial \theta'} f(x_i; \theta) + \frac{\partial \log f(x_i; \theta)}{\partial \theta} \frac{\partial f(x_i; \theta)}{\partial \theta'} dx_i \\ &= \int \frac{\partial^2 \log f(x_i; \theta)}{\partial \theta \partial \theta'} f(x_i; \theta) dx_i + \int \frac{\partial \log f(x_i; \theta)}{\partial \theta} \frac{\partial \log f(x_i; \theta)}{\partial \theta'} f(x_i; \theta) dx_i \\ &= E\left(\frac{\partial^2 \log f(X_i; \theta)}{\partial \theta \partial \theta'}\right) + E\left(\frac{\partial \log f(X_i; \theta)}{\partial \theta} \frac{\partial \log f(X_i; \theta)}{\partial \theta'}\right), \end{aligned}$$

i.e.,

$$-E\left(\frac{\partial^2 \log f(X_i; \theta)}{\partial \theta \partial \theta'}\right) = E\left(\frac{\partial \log f(X_i; \theta)}{\partial \theta} \frac{\partial \log f(X_i; \theta)}{\partial \theta'}\right) = V\left(\frac{\partial \log f(X_i; \theta)}{\partial \theta}\right) = \Sigma_i$$

Thus, $\frac{\partial \log f(X_i; \theta)}{\partial \theta}$ is distributed with mean 0 and variance Σ_i .

Note as follows:

$$I(\theta) = -E\left(\frac{\partial^2 \log L(\theta)}{\partial \theta \partial \theta'}\right) = -\sum_{i=1}^n E\left(\frac{\partial^2 \log f(X_i; \theta)}{\partial \theta \partial \theta'}\right) = \sum_{i=1}^n \Sigma_i.$$

Using the central limit theorem (generalization) shown above, asymptotically we obtain the following distribution:

$$\frac{1}{\sqrt{n}} \frac{\partial \log L(\theta)}{\partial \theta} = \frac{1}{\sqrt{n}} \sum_{i=1}^n \frac{\partial \log f(X_i; \theta)}{\partial \theta} \longrightarrow N(0, \Sigma),$$

where $\Sigma = \lim_{n \rightarrow \infty} \left(\frac{1}{n} I(\theta)\right)$.