

Review Notes for EC770

Little Tiger

4th Dec 07

Contents

Introduction	ix
I EC770 Statistics	1
1 Probability Theory	3
1.1 Lecture 1	3
Set Theory	3
Property of Set Operation	3
Collection of sets	3
1.2 Lecture 2	4
Definition of Probability	4
Property of Probability	4
Inequalities of Probability	6
Probability for limit of sequence of events	6
1.3 Lecture 3	7
Conditional probability	7
Multiplicative Rule	7
Independent Event	7
Conditional Independence	7
Law of Total Probability	7
Bayes Theorem	8
1.4 Lecture 4	8
Random Variables	8
Distributive Function of a random variable	8
Properties of c.d.f.	8
1.5 Lecture 5	9
Continuous Random Variable	9
Discrete Random Variable	9
Distribution of Two Random Variables	10
Properties of Joint Distribution Function	10
Conditional Independence	10
Independence Condition from Marginal Distribution	10
1.6 Lecture 6	11

	Joint Probability Function	11
	Joint Probability Density Function	11
	Marginal Density Fuction	12
	Marginal distribution Fuction	12
	Conditional Distribution Fuction	12
	Conditional Probability Density Fuction	12
	Generalized Multiplicative Rule	12
	Generalized Law of Total Probability	12
	Generalized Bayes Theorem	12
	Joint Distribution for k random variable	13
1.7	Lecture 7	13
	Function of Random Variable	13
	Transformation of Discrete random variable	13
	Transformation of Continuous random variable	13
	Transformation of Several random variables	14
	Useful Transformation	15
	Transformation by Convolution	16
1.8	Lecture 8	16
	Expectation of Discrete Random Variable	16
	Expectation of Continuous Random Variable	16
	Properties of Expectation	16
1.9	Lecture 9	18
	Expectation of function of random variable	18
	Moments	18
	Properties of Variance	18
	Skewness and Kurtosis	19
	Median and Mode	19
	Quartile, Interquartile range and Range	19
	Bivariate Moments	19
	Correlation	20
1.10	Lecture 10	20
	Generating Function	20
	Probability Generating Function	20
	Moment Generating Function	21
	Characteristic Function	21
	Properties of Characteristic Function	21
1.11	Lecture 11	23
	Normal Distribution	23
	Standard Normal Distribution	23
	Property of Normal distribution	23
	Multivariate normal distribution	24
	Distribution derved from Normal	24
	Other Distribution	26
1.12	Lecture 12	27
	Markov Inequality	27
	Chebyshev's Inequality	27

Generalized Markov Inequality	27
Berkstein Inequality	28
Jensen's Inequality	28
Covariance inequality	28
Point-wise Convergence	29
Norm	29
Norm Convergence	29
Convergence in Probability (Weak Convergence)	29
Almost Sure Convergence (Strong Convergence)	29
Convergence in r -th Mean	30
Convergence in Distribution (Weakest Convergence)	30
Interrationship between convergences	30
Special relationship between convergences	30
Law of Large number and Central Limit Theorem	30
2 Statistical Inference	33
2.1 Lecture 14	33
Estimation	33
Point Estimator	33
Interval Estimator	33
Unbiasedness	33
Variance of Estimator	34
Mean-square Error (MSE)	34
Consistency	34
Sufficiency	34
Factorization Theorem	34
2.2 Lecture 15	35
Point Estimation of μ	35
Point Estimation of σ^2	36
Distribution of $\bar{x}$	37
Distribution of s^2	37
2.3 Lecture 16	39
Method of Moment (MOM)	39
Generalized Method of Moment (GMM)	39
Properties of GMM Estimator	40
Maximum likelihood estimation (MLE)	41
Properties of MLE	42
MLE for normal i.i.d random sample	42
2.4 Lecture 17	42
Confidence Interval	42
Confidence interval for normal random sample	43
C.I. for GMM	44
2.5 Lecture 18	44
Hypothesis Testing	44
Testing Statistics	45
Decision error	45

Choosing testing procedure	45
Relationship to confidence Interval for normal random sample . .	45
2.6 Lecture 19	46
Power function	46
Neyman-Pearson Lemma	47
Significance level of composite null hypothesis	47
Power of composite alternative hypothesis	47
Uniformly most powerful Test (UMP Test)	47
Monotone Likelihood Ratio	48
Testing with monotone likelihood ratio	48
A The Appendix	49
Afterword	51
Acknowledgement	53
Bibliography	55

Preface

This is note for probability Theory based on classroom notes.

Introduction

This is intended to be a complete review notes including all statistics courses in Ph.D exams: Statistics, Econometrics, Time Series and Cross section Analysis. For now, only statistics is included.

Part I

EC770 Statistics

Chapter 1

Probability Theory

1.1 Lecture 1

Sample Space Ω : the collection of all possible outcome of an experiment

Event E : A collection of outcome

Experiment: a process that could be repeated

Set Theory

Terminology: empty set $\emptyset$, finite and infinite sets, disjoint sets.

Operation of Set: union $\cup$, intersection $\cap$, complement C , difference $-$, subset $\subseteq$, element of a set $\in$.

Property of Set Operation

1. Commutative: $A \cup B = B \cup A$, $A \cap B = B \cap A$
2. Associative: $A \cup (B \cup C) = (A \cup B) \cup C$, $(A \cap B) \cap C = A \cap (B \cap C)$
3. Distributive: $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$, $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$
4. DeMorgan's Law: $(A \cap B)^C = A^C \cup B^C$, $(A \cup B)^C = A^C \cap B^C$
5. Reflexive Law: $A \cup A = A$, $A \cap A = A$

Collection of sets

1. Partition: Sets A_i ($i = 1, 2, \dots, n$) where $A_i \cap A_j = \emptyset$ ($i \neq j$) are partition of S if $S = \bigcup_{i=1}^n A_i$. Note that n can be infinite.
2. Field: closed under finite union and closed under complementation. Field $\mathcal{A}$ is a collection of sets $A_1, \dots, A_n$ satisfying the following properties:

- (a) If $A_1, A_2, \dots, A_n \in \mathcal{A}$, then $\bigcup_{i=1}^n A_i \in \mathcal{A}$.
 - (b) If $A \in \mathcal{A}$, then $A^C \in \mathcal{A}$.
3. σ -field: closed under countable union and closed under complementation. Field $\mathcal{A}$ is a collection of sets $A_1, A_2, \dots$ satisfying the following properties:
- (a) If $A_1, A_2, \dots \in \mathcal{A}$, then $\bigcup_{i=1}^{\infty} A_i \in \mathcal{A}$.
 - (b) If $A \in \mathcal{A}$, then $A^C \in \mathcal{A}$.

If $\mathcal{A}_1, \mathcal{A}_2$ are σ -fields, $\mathcal{A}_1 \cap \mathcal{A}_2$ is also σ -field. In general, $\mathcal{A}_1 \cup \mathcal{A}_2$ is not necessary a σ -field. However, there exists a unique smallest σ -field that contains all elements of $\mathcal{A}_1$ and $\mathcal{A}_2$.

1.2 Lecture 2

Definition of Probability

(Kolmogorov Axiom of probability) Given a sample space Ω and an associated σ -field that contains subsets of Ω , a probability defined as P , is a function (of sets) from $\mathcal{A}$ to $[0, 1]$ satisfying the following properties:

1. $\forall A \in \mathcal{A}, P(A) \geq 0$
2. $P(\Omega) = 1$
3. Additive Property: For disjoint events $A_1, A_2, \dots \in \mathcal{A}$,

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

Therefore, probability space is defined as $(\Omega, \mathcal{A}, P)$

Property of Probability

1. If $\bigcup_{i=1}^n A_i = \Omega$ and $A_1, \dots, A_n$ is partition of Ω , then $\sum_{i=1}^n P(A_i) = 1$.
Proof: Direct application of additive property. $P(\Omega) = P\left(\bigcup_{i=1}^n A_i\right) = \sum_{i=1}^n P(A_i) = 1$
2. $P(A^C) = 1 - P(A)$
Proof: special case of property 1 when $A_1 = A$ and $A_2 = A^C$.
3. $P(\emptyset) = 0$
Proof: special case of property 2 when $A_1 = \Omega$ and $A_2 = \emptyset$

4. If $A \subset B$, then $P(A) \leq P(B)$.

Proof: From non-negativity of probability.

$$\begin{aligned} B &= B \cap \Omega = B \cap (A \cup A^C) \\ &= (B \cap A) \cup (B \cap A^C) \\ &= A \cup (B \cap A^C) \end{aligned}$$

so we have

$$\begin{aligned} P(B) &= P(A \cup (B \cap A^C)) \\ &= P(A) + P(B \cap A^C) \\ &\geq P(A) \end{aligned}$$

5. $0 \leq P(A) \leq 1$

Proof: Monotonicity of probability. Non-negativity is from $P(A) \geq 0$ and bounded by 1 is by second and third axiom.

6. $P(A \cup B) = P(A) + P(B) - P(A \cap B)$

Proof:

$$\begin{aligned} A \cup B &= A \cup (\Omega \cap B) = A \cup [(A \cup A^C) \cap B] \\ &= A \cup [(A \cap B) \cup (A^C \cap B)] \\ &= [A \cup (A \cap B)] \cup [A \cup (A^C \cap B)] \\ &= A \cup [A \cup (A^C \cap B)] = A \cup (A^C \cap B) \\ &= (A \cap \Omega) \cup (A^C \cap B) = [A \cap (B \cup B^C)] \cup (A^C \cap B) \\ &= (A \cap B) \cup (A \cap B^C) \cup (A^C \cap B) \end{aligned}$$

so that

$$P(A \cup B) = P(A \cap B) + P(A \cap B^C) + P(A^C \cap B)$$

However, we know that

$$\begin{aligned} P(A) &= P(A \cap B) + P(A \cap B^C) \\ P(B) &= P(A \cap B) + P(A^C \cap B) \end{aligned}$$

Hence,

$$\begin{aligned} P(A \cup B) &= P(A \cap B) + [P(A) - P(A \cap B)] + [P(B) - P(A \cap B)] \\ &= P(A) + P(B) - P(A \cap B) \end{aligned}$$

Inequalities of Probability

1. Bonferroni's Inequality:

$$P(A \cap B) \geq P(A) + P(B) - 1$$

Proof: $P(A \cup B) = P(A) + P(B) - P(A \cap B) \geq P(A) + P(B) - 1$ since $0 \leq P(A \cap B) \leq 1$

2. Bode's inequality:

$$P(\bigcup_{i=1}^{\infty} A_i) \leq \sum_{i=1}^{\infty} P(A_i)$$

Proof: Omitted.

3. Extension of Property 6:

$$\begin{aligned} P(A \cup B \cup C) &= P(A) + P(B) + P(C) - [P(A \cap B) + P(A \cap C) + P(B \cap C)] \\ &\quad + P(A \cap B \cap C) \end{aligned}$$

or more generally,

$$P(\bigcup_{i=1}^n A_i) = \sum_{i=1}^n P(A_i) - \sum_{i < j} P(A_i \cap A_j) + \cdots + (-1)^{q+1} \sum_{i_1 < \dots < i_q} P(\bigcap_{i=1}^q A_{i_j}) + \cdots$$

Probability for limit of sequence of events

1. Let A_i be an increasing sequence of events ($A_i \subseteq A_{i+1}$). Define $A = \lim_{i \rightarrow \infty} A_i = \bigcup_{i=1}^{\infty} A_i$. Then $P(A) = \lim_{i \rightarrow \infty} P(A_i)$.

Proof: Define $A_x = (-\infty, x]$ which is an increasing sequence set in x as $A_y \subseteq A_z$ if $y \leq z$. Hence, by property 4,

$$P(X \in A_y) \leq P(X \in A_z)$$

Define $D_k = A_k - A_{k+1}$, $A_0 = \emptyset$ and $D_1 = A_1$, then D_k is sequence of disjoint sets. By construction, we have

$$\bigcup_{i=1}^n A_i = \bigcup_{i=1}^n D_i; \quad \bigcup_{i=1}^{\infty} A_i = \bigcup_{i=1}^{\infty} D_i$$

so that

$$\begin{aligned} P\left(\lim_{i \rightarrow \infty} A_i\right) &= P(\bigcup_{i=1}^{\infty} A_i) = P(\bigcup_{i=1}^{\infty} D_i) = \sum_{i=1}^{\infty} P(D_i) \\ &= \lim_{n \rightarrow \infty} \sum_{i=1}^n P(D_i) = \lim_{n \rightarrow \infty} \sum_{i=1}^n P(A_k - A_{k+1}) \\ &= \lim_{n \rightarrow \infty} \sum_{i=1}^n [P(A_k) - P(A_{k+1})] = \lim_{n \rightarrow \infty} P(A_n) \end{aligned}$$

2. Let A_i be a decreasing sequence of events ($A_i \supseteq A_{i+1}$). Define $A = \lim_{i \rightarrow \infty} A_i = \bigcap_{i=1}^{\infty} A_i$. Then $P(A) = \lim_{i \rightarrow \infty} P(A_i)$.

Proof: Similar.

1.3 Lecture 3

Conditional probability

Definition: If $P(A) > 0$, then the conditional probability $P(B | A)$ is defined as the proportion of total probability $P(A)$ that is represented by the probability of $P(A \cap B)$. In formula,

$$P(B | A) = \frac{P(A \cap B)}{P(A)}$$

Multiplicative Rule

From the definition of conditional probability, we have

$$P(A \cap B) = P(B | A) P(A)$$

or generally, for finite k ,

$$P\left(\bigcap_{i=1}^k A_i\right) = P(A_1) P(A_2 | A_1) \cdots P\left(A_k | \bigcap_{i=1}^{k-1} A_i\right)$$

Independent Event

Event A and B are independent if

$$P(A \cap B) = P(A) P(B)$$

or if $P(A) > 0$ and $P(B) > 0$, then A and B are independent if

$$P(A | B) = P(A) \text{ and } P(B | A) = P(B)$$

For k events, $A_1, \dots, A_k$ are independent iff for all $j = 2, \dots, k$ and $(i_1, \dots, i_j) \subseteq (1, \dots, k)$

$$P(A_{i_1} \cap \cdots \cap A_{i_k}) = P(A_{i_1}) \cdots P(A_{i_j})$$

Conditional Independence

A, B, C are events in probability space with $P(A) > 0$, then B and C are conditionally independent given A iff

$$P(B \cap C | A) = P(B | A) P(C | A)$$

Law of Total Probability

Consider A_i be the partition of Ω , that is, $\sum_{i=1}^k P(A_i) = 1$, $A_i \cap A_j = \emptyset$, $\bigcup_{i=1}^k A_i = \Omega$. For any event $B \in \Omega$, we have

$$\begin{aligned} P(B) &= P\left(\bigcup_{i=1}^k (B \cap A_i)\right) = \sum_{i=1}^k P(B \cap A_i) \\ &= \sum_{i=1}^k P(A_i) P(B | A_i) \end{aligned}$$

Bayes Theorem

Let event $B_1, \dots, B_k$ be partition of Ω , $P(B_i) > 0$ and event A in the same space with $P(A) > 0$, then

$$P(B_j | A) = \frac{P(B_j) P(A | B_j)}{\sum_{i=1}^k P(B_i) P(A | B_i)}$$

1.4 Lecture 4

Random Variables

A random variable X is a function $X(\omega) : \Omega \rightarrow \mathbb{R}$ such that $\forall B \in \mathcal{B}$, $X^{-1}(B) \in \mathcal{F}$ where $\mathcal{F}$ is σ -field in original sample space and $\mathcal{B}$ is the Borel σ -field which is the smallest σ -field that contains all open intervals of $\mathbb{R}$.

Random variable transforms the probability space from $(\Omega, \mathcal{F}, P)$ to $(\mathbb{R}, \mathcal{B}, P_X)$

Distributive Function of a random variable

Also known as cumulative distribution function, or c.d.f.. In symbol, it is

$$F(x) = P_X(X \leq x) = P_X(X \in (-\infty, x]) \quad \forall x \in \mathbb{R}$$

Properties of c.d.f.

1. $0 \leq F(x) \leq 1$

Proof: $0 \leq P(x) \leq 1$

2. $F(x)$ is non-decreasing in x .

Proof: $A_x = (-\infty, x]$ is an increasing sequence set in x , so $A_x \subseteq A_y$ if $x \leq y$. Hence,

$$P(X \in A_x) \leq P(X \in A_y)$$

3. Define $F(-\infty) = \lim_{n \rightarrow \infty} F(-n)$. $F(-\infty) = 0$

Proof: Taking $A_n = (-\infty, n]$,

$$\begin{aligned} F(-\infty) &= \lim_{n \rightarrow -\infty} P(X \in A_n) = P\left(X \in \lim_{n \rightarrow -\infty} A_n\right) \\ &= P(X \in A_{-\infty}) = P(X \in \emptyset) = 0 \end{aligned}$$

4. Define $F(\infty) = \lim_{n \rightarrow \infty} F(n)$. $F(\infty) = 1$

Proof: Similar.

5. $F(x)$ is right continuous, that is

$$\lim_{\varepsilon \searrow 0} F(x + \varepsilon) = F(x)$$

Proof: let $A_n = \left(-\infty, y + \frac{1}{n}\right]$ which converge to $A_\infty = (-\infty, y]$. Then

$$\begin{aligned}\lim_{\varepsilon \rightarrow 0} F(x + \varepsilon) &= \lim_{n \rightarrow \infty} P(X \in A_n) = P\left(X \in \lim_{n \rightarrow \infty} A_n\right) \\ &= P(X \in (-\infty, y]) = F(x)\end{aligned}$$

1.5 Lecture 5

Continuous Random Variable

If $F(x)$ is differentiable, we consider the derivative of distribution function $F(x)$ and define it as the probability density function:

$$f(x) = \frac{dF(x)}{dx}$$

which implies

$$F(x) = \int_{-\infty}^x f(u) du$$

with following properties:

1. $f(x) \geq 0$

Proof: $F(x)$ is non-negative for all x .

2. $\int_{-\infty}^{\infty} f(x) dx = 1$

Proof: $F(\infty) = 1$

3. $P(a \leq x \leq b) = \int_a^b f(x) dx$

Proof: $P(x \leq b) - P(x \leq a) = F(b) - F(a) = \int_{-\infty}^b f(u) du - \int_{-\infty}^a f(u) du = \int_a^b f(x) dx$

Discrete Random Variable

Random variable X takes value on discrete points $x_1, \dots, x_n$ with $P(X = x_j) = p_j$. The probability function is $f(x_j) = p_j$ with the following properties:

1. $f(x_j) \geq 0$

2. $\sum_j f(x_j) = 1$

3. $F(x) = P(X \leq x) = \sum_{i: x_i \leq x} f(x_i)$

Distribution of Two Random Variables

Joint Distribution Function of X, Y

$$F(x, y) = P(X \leq x, Y \leq y)$$

Marginal distribution function of X and Y

$$F_X(x) = P(X \leq x) = P(X \leq x, Y < \infty) = F(x, \infty)$$

$$F_Y(y) = P(Y \leq y) = P(X < \infty, Y \leq y) = F(\infty, y)$$

Properties of Joint Distribution Function

1. $0 \leq F_X(x) \leq 1, 0 \leq F_Y(y) \leq 1$

Proof: Any probability value is bounded by 0 and 1

2. $F(x, y)$ is non-decreasing. That is, for all $a, b > 0$,

$$F(x + a, y + b) - F(x + a, y) - F(x, y + b) + F(x, y) \geq 0$$

Proof: Simple Graphic would justify it.

3. $F(-\infty, y) = 0$ and $F(x, -\infty) = 0$

4. $F(\infty, \infty) = 1$

5. $F(x, y)$ is right continuous, that is,

$$\lim_{\varepsilon_1 \downarrow 0, \varepsilon_2 \downarrow 0} F(x + \varepsilon_1, y + \varepsilon_2) = \lim_{\varepsilon_1 \downarrow 0} F(x + \varepsilon_1, y) = \lim_{\varepsilon_2 \downarrow 0} F(x, y + \varepsilon_2) = F(x, y)$$

Conditional Independence

Events A and B are independent if for all Borel set $A, B \in \mathcal{B}$,

$$P(X \in A \mid Y \in B) = \frac{P(X \in A, Y \in B)}{P(Y \in B)}$$

or

$$\begin{aligned} P_{X|Y}(A \mid B) &= P_X(A) \quad \forall A, B \in \mathcal{B} \\ P_{X \cap Y}(A \cap B) &= P_X(A) P_Y(B) \quad \forall A, B \in \mathcal{B} \end{aligned}$$

Independence Condition from Marginal Distribution

Given existence of $F(x, y), F_X(x), F_Y(y)$, we have

$$F(x, y) = F_X(x) F_Y(y) \quad \forall x, y$$

if and only if X and Y are independent.

Proof: "Only if": We need to show for all $(a_1, b_1]$ and $(a_2, b_2]$, we have

$$P((a_1, b_1] \cap (a_2, b_2]) = P_X((a_1, b_1]) \cdot P_Y((a_2, b_2])$$

Now, we have

$$\begin{aligned} P((a_1, b_1] \cap (a_2, b_2]) &= P(a_1 < x \leq b_1, a_2 < y \leq b_2) \\ &= P(X \leq b_1, Y \leq b_2) - P(X < a_1, Y \leq b_2) - P(X \leq b_1, Y < a_2) + P(X \leq a_1, Y \leq a_2) \\ &= F(b_1, b_2) - F(a_1, b_2) - F(b_1, a_2) + F(a_1, a_2) \\ &= F_X(b_1) F_Y(b_2) - F_X(a_1) F_Y(b_2) - F_X(b_1) F_Y(a_2) + F_X(a_1) F_Y(a_2) \\ &= F_X(b_1) [F_Y(b_2) - F_Y(a_2)] - F_X(a_1) [F_Y(b_2) - F_Y(a_2)] \\ &= [F_X(b_1) - F_X(a_1)] [F_Y(b_2) - F_Y(a_2)] \\ &= P_X((a_1, b_1]) \cdot P_Y((a_2, b_2]) \end{aligned}$$

"If": Since X and Y are independent, taking $A = (-\infty, x], B = (-\infty, y]$, we have

$$P((a_1, b_1] \cap (a_2, b_2]) = P_X((a_1, b_1]) \cdot P_Y((a_2, b_2])$$

so that

$$F(x, y) = F_X(x) F_Y(y) \quad \forall x, y$$

1.6 Lecture 6

Joint Probability Function

Joint probability function is denoted as $f(x_i, y_j) = P(X = x_i, Y = y_j) = p_{ij}$

Marginal probability function for X : $P(X = x_i) = \sum_j P(X = x_i, Y = y_j) = \sum_j p_{ij} = p_{i\cdot}$

Marginal probability function for Y : $P(Y = y_j) = \sum_i P(X = x_i, Y = y_j) = \sum_i p_{ij} = p_{\cdot j}$

Conditional probability function:

$$P(X = x_i | Y = y_j) = \frac{P(X = x_i, Y = y_j)}{P(Y = y_j)} = \frac{p_{ij}}{p_{\cdot j}}$$

Joint Probability Density Function

If $F(x)$ is differentiable, then

$$\frac{\partial^2 F(x, y)}{\partial x \partial y} = f(x, y)$$

is called joint density function with following properties (state without proof):

1. $f(x, y) \geq 0$
2. $\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(x, y) dx dy = 1$
3. $\int_{-\infty}^x \int_{-\infty}^y f(u, v) du dv = F(x, y)$

Marginal Density Fuction

$$f_X(x) = \int_{-\infty}^{\infty} f(x, y) dy; \quad f_Y(y) = \int_{-\infty}^{\infty} f(x, y) dx$$

Marginal distribution Fuction

$$F_X(x) = \int_{-\infty}^x f_X(u) du; \quad F_Y(y) = \int_{-\infty}^y f_Y(v) dv$$

Conditional Distribution Fuction

$$\begin{aligned} P(X \leq x \mid Y = y) &= F(x \mid y) \\ &= \lim_{\delta \rightarrow 0} P(X \leq x \mid y - \delta \leq Y \leq y + \delta) \\ &= \lim_{\delta \rightarrow 0} \frac{P(X \leq x, y - \delta \leq Y \leq y + \delta)}{P(y - \delta \leq Y \leq y + \delta)} \frac{2\delta}{2\delta} \\ &= \lim_{\delta \rightarrow 0} \frac{\int_{-\infty}^x \int_{y-\delta}^{y+\delta} f(u, v) dv du}{\int_{y-\delta}^{y+\delta} f_Y(v) dv} \frac{2\delta}{2\delta} \\ &= \lim_{\delta \rightarrow 0} \frac{\int_{-\infty}^x \frac{F(u, y + \delta) - F(u, y - \delta)}{2\delta} du}{\frac{F_Y(y + \delta) - F_Y(y - \delta)}{2\delta}} \\ &= \frac{\int_{-\infty}^x f(u, y) du}{f_Y(y)} \end{aligned}$$

Conditional Probability Density Fuction

$$f(x \mid y) = \frac{\partial F(x, y)}{\partial y} = \frac{f(x, y)}{f_Y(y)}$$

Generalized Multiplicative Rule

$$f(x, y) = f_X(x) f_Y(y \mid x) = f_Y(y) f_X(x \mid y)$$

Generalized Law of Total Probability

$$f_X(x) = \int_{-\infty}^{\infty} f(x, y) dy = \int_{-\infty}^{\infty} f_X(x \mid y) f_Y(y) dy$$

Generalized Bayes Theorem

$$f(x \mid y) = \frac{f_X(x) f_Y(y \mid x)}{\int_{-\infty}^{\infty} f_X(x \mid y) f_Y(y) dy}$$

Joint Distribution for k random variable

$$F(X_1, \dots, X_k) = P(X_1 \leq x_1, \dots, X_k \leq x_k)$$

For discrete case, probability distribution is completely determined by

$$P(X_1 = x_1, \dots, X_k = x_k)$$

For continuous case,

$$f(x_1, \dots, x_k) = \frac{\partial^k F(x_1, \dots, x_k)}{\partial x_1 \cdots \partial x_k}$$

1.7 Lecture 7**Function of Random Variable**

A function of a random variable $g(X)$ is also a random variable as long as the mapping is measurable, that is,

$$\forall B \in \mathcal{B}, \quad P(g(X) \in B)$$

is defined if $X^{-1}(g^{-1}(B)) \in \mathcal{F}$.

Transformation of Discrete random variable

Discrete Case: For $Y = r(X)$,

$$P(Y = y) = P(r(X) = y) = \sum_{i: r(x_i) = y} P(X = x_i)$$

Transformation of Continuous random variable

Method I (Inversion method): For $Y = r(X)$,

$$P(Y \leq y) = P(r(X) \leq y) = \int_{x: r(x) \leq y} f(x) dx$$

Method II (Jacobian transformation):

Thm: Suppose $P(a < X < b) = 1$, $Y = r(X)$ is a continuous and monotone function for $a < X < b$. When $a < X < b$, $\alpha < Y < \beta$. Then the p.d.f. of Y is given by

$$g(y) = f(r^{-1}(y)) \left| \frac{dx}{dy} \right|$$

for all $\alpha < y < \beta$ where $f(\cdot)$ is pdf of X . In particular,

if $r(\cdot)$ is an increasing function, $g(y) = f(r^{-1}(y)) \frac{dx}{dy}$.

if $r(\cdot)$ is a decreasing function, $g(y) = -f(r^{-1}(y)) \frac{dx}{dy}$.

Proof: If r is increasing, then

$$G(y) = P(Y \leq y) = P(r(X) \leq y) = P(X \leq r^{-1}(y)) = F(r^{-1}(y))$$

which implies

$$g(y) = f(r^{-1}(y)) \frac{dx}{dy}$$

Similarly, if r is decreasing,

$$G(y) = P(Y \leq y) = P(r(X) \leq y) = P(X \geq r^{-1}(y)) = 1 - F(r^{-1}(y))$$

which implies

$$g(y) = -f(r^{-1}(y)) \frac{dx}{dy}$$

Transformation of Several random variables

$X_1, \dots, X_n$ are random variables and the new random variables $Y_1, \dots, Y_m$ are represented by function

$$\begin{cases} Y_1 = r_1(X_1, \dots, X_n) \\ \vdots \\ Y_m = r_m(X_1, \dots, X_n) \end{cases}$$

Discrete Case:

$$P(Y_1 = y_1, \dots, Y_n = y_n) = \sum_{x=(x_1, \dots, x_n); r_j(x)=y_j} P(X = x_j)$$

Continuous Case:

Method I (Inversion method)

$$\begin{aligned} G(y_1, \dots, y_m) &= P(Y_1 \leq y_1, \dots, Y_m \leq y_m) \\ &= \int \cdots \int_{r_j(x_1, \dots, x_n) \leq y_j; j=1, \dots, m} f(x_1, \dots, x_n) dx_1 \cdots dx_n \end{aligned}$$

Method II (Jacobian matrix)

For case of one-to-one mapping ($m = n$)

If $P((x_1, \dots, x_n) \in S) = 1$ and there is a one-to-one correspondence between $(X_1, \dots, X_n)$ and $(Y_1, \dots, Y_n)$, $Y_j = r_j(X_1, \dots, X_n)$, then the joint p.d.f. of $(Y_1, \dots, Y_n)$ can be obtained from

$$g(y_1, \dots, y_n) = f(x_1(y), \dots, x_n(y)) \left| \frac{dx}{dy} \right|$$

where

$$\left| \frac{dx}{dy} \right| = \begin{vmatrix} \frac{dx_1}{dy_1} & \cdots & \frac{dx_1}{dy_n} \\ \vdots & \ddots & \vdots \\ \frac{dx_n}{dy_1} & \cdots & \frac{dx_n}{dy_n} \end{vmatrix}$$

is called the Jacobian of Transformation.

Speical Case I: r is linear combination of $X_1, \dots, X_n$.

$$\begin{bmatrix} Y_1 \\ \vdots \\ Y_n \end{bmatrix} = Y = AX = A \begin{bmatrix} X_1 \\ \vdots \\ X_n \end{bmatrix}$$

and A is non-singular, then we have

$$X = A^{-1}Y; \quad \left| \frac{dx}{dy} \right| = |A^{-1}|$$

so that

$$g(y) = \frac{1}{|A|} f(A^{-1}y)$$

Useful Transformation

$X_1, \dots, X_n \sim F(\cdot)$, $f(x_i)$ are iid, $Y_n = \max\{X_1, \dots, X_n\}$ and $Z_n = \min\{X_1, \dots, X_n\}$.
Now,

$$\begin{aligned} P(Y_n \leq y) &= P(\max\{X_1, \dots, X_n\} \leq y) \\ &= P(X_1 \leq y, \dots, X_n \leq y) \\ &= P(X_1 \leq y) \dots P(X_n \leq y) \\ &= [F(y)]^n \equiv G(y) \end{aligned}$$

and if X_i are continuous,

$$g(y) = n [F(y)]^{n-1} f(y)$$

Now,

$$\begin{aligned} P(Z_n \leq z) &= P(\min\{X_1, \dots, X_n\} \leq z) \\ &= 1 - P(\min\{X_1, \dots, X_n\} \geq z) \\ &= 1 - P(X_1 \geq z, \dots, X_n \geq z) \\ &= 1 - [P(X_1 \geq z) \dots P(X_n \geq z)] \\ &= 1 - [1 - F(z)]^n \equiv H(y) \end{aligned}$$

and if X_i are continuous,

$$h(z) = \frac{dH(z)}{dz} = n [1 - F(y)]^{n-1} f(y)$$

so

$$\begin{aligned} P(y \leq X_1, \dots, X_n \leq z) &= F(Y_n \leq y, Z_n \leq z) \\ &= P(Y_n \leq y) - P(Y_n \leq y, Z_n > z) \\ &= P(Y_n \leq y) - P(z < X_1 \leq y, \dots, z < X_n \leq y) \\ &= [F(y)]^n - [F(y) - F(z)]^n \end{aligned}$$

Transformation by Convolution

X_1 and X_2 are two random variables and let $Y = X_1 + X_2$ and $Z = X_2$ so that

$$\begin{bmatrix} Y \\ Z \end{bmatrix} = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} X_1 \\ X_2 \end{bmatrix} = AX \quad \text{where } |A^{-1}| = 1$$

where

$$\begin{aligned} g(y, z) &= f(x_1(\cdot), x_2(\cdot)) \frac{1}{|A|} \\ &= f(y - z, z) \end{aligned}$$

so that

$$g(y) = \int f(y - z, z) dz$$

Furthermore, if X_1 and X_2 are independent, then

$$g(y) = \int f_1(y - z) f_2(z) dz$$

1.8 Lecture 8

Expectation of Discrete Random Variable

For a discrete random variable X taking values $\{x_j\}_{j=1}^J$, with probability function $P(X = x_j) = p_j$, the expectation of X is defined as

$$EX = \sum_{j=1}^J x_j p_j$$

Rmk: For $EX < \infty$, we need to have $\sum_{j=1}^J |x_j| p_j < \infty$.

Expectation of Continuous Random Variable

The expectation of continuous random variable X with pdf $f(x)$ is

$$EX = \int_{-\infty}^{\infty} x f(x) dx$$

Rmk 1: EX exists if $\int_{-\infty}^{\infty} |x| f(x) dx < \infty$.

Rmk 2: Since $dF(X) = f(x) dx$, we have $EX = \int_{-\infty}^{\infty} x dF(x)$.

Properties of Expectation

Rmk: For the following proof, we only show the continuous case as discrete case is similar.

1. If $P(X = c) = 1$, then $EX = c$.

Proof: Immediate from definition.

2. If EX exists, then $E(cX) = cX$.

Proof: As constant can be taken out from summation/integral.

$$\int_{-\infty}^{\infty} cf(x) dx = c \int_{-\infty}^{\infty} f(x) dx = c$$

3. If $Y = aX + b$ and EX exists, $EY = aEX + b$.

Proof: EY exists because

$$|Y| \leq |a||X| + |b| \Rightarrow \int_{-\infty}^{\infty} |y| dF(y) \leq |a| \int_{-\infty}^{\infty} |x| dF(x) + |b| < \infty$$

4. If $Y = X_1 + X_2$, EX_1 and EX_2 exist, then $EY = EX_1 + EX_2$.

Proof: EY exists because

$$\begin{aligned} & \iint |x_1 + x_2| f(x_1, x_2) dx_1 dx_2 \\ & \leq \iint (|x_1| + |x_2|) f(x_1, x_2) dx_1 dx_2 \\ & = \iint |x_1| f(x_1, x_2) dx_1 dx_2 + \iint |x_2| f(x_1, x_2) dx_1 dx_2 \\ & = \int |x_1| f(x_1) dx_1 + \int |x_2| f(x_2) dx_2 < \infty \end{aligned}$$

where the result is just additive property of integration.

5. Linear property: If EX_i exists, $E(a_1X_1 + \dots + a_kX_k + b) = a_1EX_1 + \dots + a_kEX_k + b$

Proof: General case of property 4.

6. If $P(X \geq a) = 1$, a is a constant, then $EX \geq a$.

Proof:

$$\begin{aligned} EX &= \int x dF(x) = \int_{x < a} x dF(x) + \int_{x \geq a} x dF(x) \\ &= \int_{x \geq a} x dF(x) \geq \int_{x \geq a} a dF(x) = aP(X \geq a) = a \end{aligned}$$

7. If $P(X \geq a) = 1$ and $EX = a$, then $P(X > a) = 0$ and $P(X = a) = 1$.

Proof:

$$EX = \int x dF(x) = \int_{x < a} x dF(x) + \int_{x \geq a} x dF(x)$$

which implies

$$a = \lim_{\varepsilon \rightarrow 0} \int_a^{a+\varepsilon} x dF(x) + \int_{a+\varepsilon}^{\infty} x dF(x)$$

and this can only be true if $P(X > a) = 0$ and $P(X = a) = 1$.

1.9 Lecture 9

Expectation of function of random variable

$$E[g(x)] = \int_{-\infty}^{\infty} g(x) dF(x)$$

Moments

Special Case I: $g(X) = X^r$, we call EX^r the r -th raw moment of X .

Thm: If EX^r exists, then for any $0 < s < r$, $EX^s < \infty$.

Proof:

$$\begin{aligned} \int |x|^s dF(x) &= \int_{|x| < 1} |x|^s dF(x) + \int_{|x| \geq 1} |x|^s dF(x) \\ &\leq \int_{|x| < 1} dF(x) + \int_{|x| \geq 1} |x|^s dF(x) \\ &\leq \int_{|x| < 1} dF(x) + \int_{|x| \geq 1} |x|^r dF(x) < \infty \end{aligned}$$

Rmk 1: If k^{th} moment does not exist, $(k+1)$ th moment will also not exist.

Rmk 2: When $r = 1$, we call EX as mean of distribution.

Special Case II: $g(X) = (X - EX)^r$, we call $E(X - EX)^r$ the r -th central moment of X .

Rmk 1: EX^r exists then $E(X - EX)^r$ also exists.

Rmk 2: $E(X - EX)^r = \binom{r}{0} EX^r - \binom{r}{1} EX^{r-1} EX + \binom{r}{2} EX^{r-2} EX^2 + \dots + (-1)^r \binom{r}{r} EX^r$.

Rmk 3: When $r = 2$, we call $E(X - EX)^2$ the variance of distribution.

Properties of Variance

1. $Y = aX + b$, then $\text{Var}Y = a^2 \text{Var}X$

Proof: From definition

2. If $\text{Var}X = 0$, then there exists c such that $P(X = c) = 1$

Proof: $\text{Var}(X) = 0$ implies $E(X - EX)^2 = 0$. Then this means for all X , $X = EX$, so that all X would always be equal to one constant, say c .

Skewness and Kurtosis

Skewness: Measure the symmetry of the distribution:

$$\beta_1 = \frac{E(X - EX)^3}{\sigma^3}$$

Kurtosis: Measure the tail of distribution:

$$\beta_2 = \frac{E(X - EX)^4}{\sigma^4}$$

Median and Mode

Median m is defined as $P(X \leq m) = P(X \geq m) = 0.5$.

For continuous variable,

$$m = F^{-1}(0.5)$$

where F^{-1} is inverse function of distribution function F .

Mode μ_0 is defined as $f(\mu_0) = \sup_X f(X)$.

Quartile, Interquartile range and Range

If $F(x)$ is strictly monotone, $F^{-1}(\tau)$ is called τ -th quartile of X where $\tau \in [0, 1]$

We call $Q(\tau) = F^{-1}(\tau)$ as quartile function

In general, the τ -th quartile of a distribution is defined as

$$Q_x(\tau) = \inf \{x : F(x) \geq \tau\}$$

when $\tau = 0.5$, $Q_x(\tau)$ is the median.

when $\tau = 0.25, 0.5, 0.75$, $Q_x(\tau)$ are quartiles.

when $\tau = 0.2, 0.4, 0.6, 0.8$, $Q_x(\tau)$ are quintiles.

when $\tau = 0.1, 0.2, \dots, 0.9$, $Q_x(\tau)$ are deciles.

Interquartile range IQR is defined as $IQR = Q_x(0.75) - Q_x(0.25)$

Range is defined as $range = \sup x - \inf x = Q_x(1) - Q_x(0)$.

Bivariate Moments

1. $EXY = \iint xyf(x, y) dx dy = \iint xy dF(x, y)$
2. $Eg(X, Y) = \int g(x, y) dF(x, y)$
3. $E(X^r Y^s)$ is called the product of raw moment of order r and s .
4. $E[(X - EX)^r (Y - EY)^s]$ is called the product of central moment of order r and s .
5. Particularly, when $r = 1$, $s = 1$, we have

$$\begin{aligned} Cov(X, Y) &= E[(X - EX)(Y - EY)] \\ &= EXY - EXEY \end{aligned}$$

6. Furthermore, if X and Y are independent, then $Cov(X, Y) = 0$.

Proof: $Cov(X, Y) = EXY - EXEY = EXEY - EXEY = 0$.

7. If $cov(X, Y) = 0$, then X and Y are uncorrelated.

Correlation

$$corr(X, Y) = \frac{cov(X, Y)}{\sqrt{VarX}\sqrt{VarY}}$$

Thm: Correlation coefficient is bounded by 1, $|corr(X, Y)| \leq 1$.

Proof: By Cauchy-Scharwz inequality,

$$[E(AB)]^2 \leq EA^2EB^2$$

Taking $A = X - EX$ and $B = Y - EY$, we have

$$E[(X - EX)(Y - EY)]^2 \leq E(X - EX)^2 E(Y - EY)^2$$

or

$$[cov(X, Y)]^2 \leq VarX \cdot VarY$$

so that

$$\frac{[cov(X, Y)]^2}{VarX \cdot VarY} \leq 1$$

which implies

$$|corr(X, Y)| \leq 1$$

1.10 Lecture 10

Generating Function

Given a random variable X , generating function $G_X(t)$ is defined as

$$G_X(t) = Eg(x, t) = \int g(x, t) dF(x)$$

Probability Generating Function

$$\begin{aligned} g(x, t) &= t^x \\ P_X(t) &= \int t^x dF(x) = Et^x \end{aligned}$$

Moment Generating Function

$$g(x, t) = e^{tx}$$

$$M_X(t) = Ee^{tx} = \int e^{tx} dF(x)$$

For small value of t ,

$$e^{tx} = 1 + tx + \cdots + \frac{(tx)^j}{j!}$$

so that

$$M_X(t) = \int \sum_{j=0}^{\infty} \frac{(tx)^j}{j!} dF(x) = \sum_{j=0}^{\infty} \frac{t^j}{j!} \int x^j dF(x)$$

which implies

$$\left. \frac{\partial^j M_X(t)}{\partial t^j} \right|_{t=0} = EX^j$$

Characteristic Function

$$g(x, t) = e^{itx}$$

$$\phi_X(t) = \int e^{itx} dF(x) = Ee^{itx}$$

Rmk 1: Characteristic Function is always well-defined.

$$\int e^{itx} dF(x) < \infty$$

since

$$e^{itx} = \cos tx + i \sin tx$$

so that

$$|e^{itx}| = 1.$$

Rmk 2: If X is vector of random variable, the c.f. would be, for example, say $X = (Y, Z)$

$$\phi_X(t, s) = Ee^{i(tY+sZ)} = E(e^{itY} e^{isZ})$$

Properties of Characteristic Function

1. There is one to one correspondence between c.f. and p.d.f..
2. If the r -th moment exists, then

$$EX^r = \frac{1}{i^r} \left. \frac{\partial \phi_X(t)}{\partial t^r} \right|_{t=0}.$$

Proof:

$$\begin{aligned}\frac{\partial \phi_X(t)}{\partial t^r} &= i^r \int x^r e^{itx} dF(x) \\ &= i^r \int x^r dF(x) \text{ as } t \rightarrow \infty \\ &= i^r EX\end{aligned}$$

3. Relationship between m.g.f. and c.f.:

$$\phi_X(t) = M_X(it)$$

Proof: From definition, $M_X(t) = Ee^{tX}$ and $\phi_X(t) = Ee^{itX}$

4. Linear function of random variable:

$$\phi_{a+bX}(t) = e^{ita} \phi_X(bt)$$

Proof: From definition, $\phi_{a+bX}(t) = Ee^{it(a+bX)} = e^{ita} Ee^{i(tb)X} = e^{ita} \phi_X(bt)$

5. If X and Y are independent random variables, we have

$$\phi_{X+Y}(t, s) = \phi_X(t) \phi_Y(s)$$

6. If $\left. \frac{\partial \phi_X(t)}{\partial t^r} \right|_{t=0}$ exists, then

$$\begin{cases} EX^r < \infty & \text{if } r \text{ is even} \\ EX^{r-1} < \infty & \text{if } r \text{ is odd} \end{cases}$$

7. If $|X^r| < \infty$, then

$$\phi_X(t) = \sum_{j=0}^k \frac{EX^j}{j!} (it)^j + o(t^k)$$

where $o(t^k)$ is defined as

$$\lim_{t \rightarrow \infty} \frac{o(t^k)}{t^k} = 0.$$

8. The density function $f(x)$ can be obtained by inverse transformation

$$f(x) = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{-itx} \phi_X(t) dt.$$

1.11 Lecture 11

Normal Distribution

The density function is

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

for all $x \in \mathbb{R}$, $-\infty < \mu < \infty$, $\sigma > 0$ with the mean and variance equal to

$$EX = \mu; \quad \text{Var}X = E(x - \mu)^2 = \sigma^2.$$

The r th central moment is

$$\begin{aligned} E(X - \mu)^r &= \begin{cases} 0 & \text{if } r \text{ is odd} \\ (2k-1)(2k-3)\cdots 3 \cdot 1 \cdot \sigma^{2k} & \text{if } r \text{ is even and } r = 2k \end{cases} \\ &= c_r \sigma^r \text{ where } c_r = \begin{cases} 0 & \text{if } r \text{ is odd} \\ (2k-1)(2k-3)\cdots 3 \cdot 1 & \text{if } r \text{ is even and } r = 2k \end{cases} \end{aligned}$$

The mgf and cf are

$$M_X(t) = \exp\left\{\mu t + \frac{\sigma^2 t^2}{2}\right\}; \quad \phi_X(t) = \exp\left\{i\mu t - \frac{\sigma^2 t^2}{2}\right\}$$

Standard Normal Distribution

When $\mu = 0$ and $\sigma^2 = 1$, then we call the distribution to be standard normal and the p.d.f. and c.d.f. are denoted by $\phi(x)$ and $\Phi(x)$ respectively.

Property of Normal distribution

1. If $X \sim N(\mu, \sigma^2)$, then $Y = aX + b \sim N(a\mu + b, a^2\sigma^2)$

$$\text{Proof: } \phi_Y(t) = \phi_{aX+b}(t) = e^{itb} \phi_X(at) = e^{itb} \left[e^{i\mu(at) - \sigma^2(at)^2/2} \right] = e^{i(a\mu+b)t - (a\sigma)^2 t^2/2}$$

2. If $X_1, \dots, X_k \sim N(\mu_i, \sigma_i^2)$ are i.i.d., then

$$\sum_{i=1}^k X_i \sim N\left(\sum_{i=1}^k \mu_i, \sum_{i=1}^k \sigma_i^2\right).$$

Proof: Let $Y = \sum x_i$ for $i = 1, \dots, k$ of $\phi_{X_i}(t) = \exp\left\{iu_i t - \frac{\sigma_i^2 t^2}{2}\right\}$. Since x_i are independent,

$$\begin{aligned} \phi_Y(t) &= \prod_{i=1}^k \phi_{X_i}(t) = \prod_{i=1}^k \exp\left\{iu_i t - \frac{\sigma_i^2 t^2}{2}\right\} \\ &= \exp\left\{it \sum_{i=1}^k u_i - \frac{t^2 \sum_{i=1}^k \sigma_i^2}{2}\right\}. \end{aligned}$$

3. If $Y = a + \sum_{i=1}^k b_i x_i$, then $Y \sim N\left(a + \sum_{i=1}^k b_i \mu_i, \sum_{i=1}^k b_i^2 \sigma_i^2\right)$

Proof: Combination of 1 and 2 proof.

Multivariate normal distribution

X_i ($i = 1, \dots, p$) are $N(\mu_i, \sigma_i^2)$ with covariance matrix $\Sigma_{p \times p}$ where

$$\Sigma = \begin{bmatrix} \sigma_{11} & \sigma_{12} & \cdots & \sigma_{1p} \\ \sigma_{21} & \ddots & & \vdots \\ \vdots & & \ddots & \vdots \\ \sigma_{p1} & \cdots & \cdots & \sigma_{pp} \end{bmatrix}$$

such that $\sigma_{ii} = \sigma_i^2$ and $\sigma_{ij} = \text{cov}(X_i, X_j)$.

Let $X = (X_1, X_2, \dots, X_p)^T$ is a p -dimensional normal distribution with mean $\mu = (\mu_1, \mu_2, \dots, \mu_p)^T$ and covariance matrix Σ . The joint density of X is

$$f(x) = \frac{1}{(2\pi)^{p/2} |\Sigma|^{1/2}} \exp \left\{ -\frac{1}{2} (x - \mu)^T \Sigma^{-1} (x - \mu) \right\} \quad \text{for } x \in \mathbb{R}^p$$

and we denote this by $X \sim N(\mu, \Sigma)$.

Prop: If $X \sim N(\mu, \Sigma)$, then $AX \sim N(A\mu, A\Sigma A^T)$

Proof: Omitted.

Distribution derived from Normal

Lognormal Distribution

X is lognormal if $Y = \log X \sim N(\mu, \sigma^2)$

c.f. is $\phi_X(t) = \exp \{i\mu t - \sigma^2 t^2/2\}$

m.g.f. $M_X(t) = EX^t = \exp \{\mu t + \sigma^2 t^2/2\}$

$EX = \exp(\mu + \sigma^2/2)$ and $VarX = e^{2\mu + \sigma^2} (e^{\sigma^2} - 1)$

χ^2 Distribution

If X_i ($i = 1, \dots, p$) are independent $N(0, 1)$, then $Y = \sum_{i=1}^p X_i^2$ is called a random variable of χ^2 distribution with p degree of freedom, which is denoted as $Y \sim \chi_p^2$.

The following are properties of chi-square distribution:

1. If $Z_i \sim N(\mu_i, \sigma_i^2)$ are independent, then

$$\sum_{i=1}^p \left(\frac{Z_i - \mu_i}{\sigma_i} \right)^2 \sim \chi_p^2$$

2. $EX = p$ and $VarX = 2p$. Other higher moments can be found using the formula:

$$E(X - \mu)^r = c_r \sigma^r \text{ where } c_r = \begin{cases} 0 & \text{if } r \text{ is odd} \\ (2k-1)(2k-3)\cdots 3 \cdot 1 & \text{if } r \text{ is even and } r = 2k \end{cases}$$

3. If $Y_1 \sim \chi^2(p)$ and $Y_2 \sim \chi^2(q)$, Y_1 and Y_2 are independent, then $Y_1 + Y_2 \sim \chi^2(p+q)$.

Non-central χ^2 - distribution

If X_i are independent $N(\mu_i, 1)$ and $\mu_i \neq 0$, then

$$Y = \sum_{i=1}^p X_i^2 \sim \chi_p^2 \left(\sum_{i=1}^p \mu_i^2 \right)$$

where $\sum_{i=1}^p \mu_i^2$ is called non-central parameter.

Properties: For $Y \sim \chi_p^2(\lambda)$, $EX = p + \lambda$ and $VarX = 2p + 4\lambda$

Proof:

$$\begin{aligned} Y &= \sum X_i^2 = \sum (X_i - \mu_i + \mu_i)^2 \\ &= \sum \left[(X_i - \mu_i)^2 + \mu_i^2 + 2(X_i - \mu_i)\mu_i \right] \\ &= \sum (X_i - \mu_i)^2 + \sum \mu_i^2 \end{aligned}$$

Student's t distribution

If $X \sim N(0, 1)$, $Y \sim \chi_p^2$ and X, Y are independent, then

$$t = \frac{X}{\sqrt{Y/p}}$$

has t distribution with degree of freedom p .

Rmk1: $f(x) = c_p \left(1 + \frac{x^2}{p}\right)^{-\frac{p+1}{2}}$ where c is a constant depend on p

Rmk2: $t_p \rightarrow N(0, 1)$ as $p \rightarrow \infty$.

Rmk3: For $X \sim t_p$ having moments up to order $(p-1)$,

1. EX^r does not exist when $r \geq p$.

2. For $r < p$,

$$EX^r = \frac{p^{r/2} \Gamma\left(\frac{r+1}{2}\right) \Gamma\left(\frac{p-r}{2}\right)}{\Gamma\left(\frac{1}{2}\right) \Gamma\left(\frac{p}{2}\right)}$$

where

$$\Gamma(a) = \int_0^\infty x^{a-1} e^{-x} dx.$$

3. $EX = 0$ for $p > 1$ and $VarX = \frac{p}{p-2}$ for $p > 2$.

F-distribution

If $Y_1 \sim \chi_m^2$, $Y_2 \sim \chi_n^2$ are independent, then

$$F_{m,n} = \frac{Y_1/m}{Y_2/n}$$

Rmk1: For $r < n/2$, EX^r exists.

Rmk2: For $n > 2$, $EX = \frac{n}{n-2}$; For $n > 4$, $VarX = \frac{2n^2(m+n-2)}{m(n-2)^2(n-4)}$; For $n \rightarrow \infty$, $EX = 1$

Rmk3: $t_q^2 = \frac{\chi_1^2}{\chi_q^2/q} \sim F(1, q)$; This can be seen by $Var(t_q) = \frac{q}{q-2} \rightarrow 1$ as $q \rightarrow \infty$ and $E(t_q) = 0$.

Other Distribution**Gamma Distribution**

$$\begin{aligned} f(x) &= \frac{1}{\Gamma(t)} \lambda^t x^{t-1} e^{-\lambda x}; x \geq 0, t > 0 \\ EX &= \frac{t}{\lambda}; \quad VarX = \frac{t}{\lambda^2} \end{aligned}$$

Bernoulli Distribution

$$X = \begin{cases} 0 & \text{with probability } p \\ 1 & \text{with probability } 1-p \end{cases}$$

so that

$$p(x) = p^x (1-p)^{1-x}$$

Binomial Distribution

If $Y = X_1 + X_2 + \dots + X_n$ where X_i are i.i.d. Bernoulli distribution, then $Y \sim B(n, p)$

Poisson Distribution

X takes the value from $0, 1, 2, \dots$. The probability function is

$$P(X = x) = \frac{e^{-\lambda} \lambda^x}{x!}$$

Indicator Function

For event $A \in \Omega$, I_A is defined as

$$I_A = \begin{cases} 0 & \omega \notin A \\ 1 & \omega \in A \end{cases}$$

Rmk: This is special case of Bernoulli.

1.12 Lecture 12

Markov Inequality

Let $g(\cdot)$ be a non-negative function of a random variable X and $E[g(X)]$ exists, then for any $c > 0$, we have

$$P[g(X) \geq c] \leq \frac{E[g(X)]}{c}$$

Proof:

$$\begin{aligned} E[g(X)] &= \int g(x) dF(x) = \int_{g(x) \geq c} g(x) dF(x) + \int_{g(x) < c} g(x) dF(x) \\ &\geq \int_{g(x) \geq c} g(x) dF(x) \geq \int_{g(x) \geq c} c dF(x) = cP[g(X) \geq c] \end{aligned}$$

If X is non-negative, that is $P(X \geq 0) = 1$. Then, let $g(X) = X$, we have

$$P(X \geq c) \leq \frac{EX}{c}.$$

In general, consider $g(x) = |x|$, then

$$P(|X| \geq c) \leq \frac{E|X|}{c}$$

Chebyshev's Inequality

Consider $g(x) = (x - \mu)^2$ where $\mu = EX$, then we have

$$P[(x - \mu)^2 \geq c] \leq \frac{E[(x - \mu)^2]}{c} = \frac{Var X}{c}$$

Furthermore, taking $c = \varepsilon^2 \sigma^2$ then,

$$P[(x - \mu)^2 \geq \varepsilon^2 \sigma^2] \leq \frac{\sigma^2}{\varepsilon^2 \sigma^2} = \frac{1}{\varepsilon^2}$$

which implies

$$P[|x - \mu| \geq \varepsilon \sigma] \leq \frac{1}{\varepsilon^2}$$

Generalized Markov Inequality

Let $g(\cdot)$ be a non-negative function on $\mathbb{R}$ and $g(\cdot)$ is non-decreasing on the range of a random variable z . Then

$$P(z \geq a) \leq \frac{E[g(z)]}{g(a)}$$

Proof:

$$P(z \geq a) = P(g(z) \geq g(a)) \leq \frac{E[g(z)]}{g(a)}$$

where the last step is application of Markov inequality.

Berkstein Inequality

Let $g(t) = e^{bt}$ for $b > 0$, then

$$P(z \geq a) \leq e^{-ab} E e^{bz} \leq \inf_{b \geq 0} e^{-ab} E e^{bz}$$

Corollary 1: If $X \sim \text{Binomial}(n, p)$, then $P(|x - np| \geq n\varepsilon) \leq 2e^{-n\varepsilon^2/2}$

Corollary 2: If u_i are independent r.v. and $P(u_i \leq b) = 1$, $v = \sum_{i=1}^n E(u_i^2)$, then for all $a > 0$, we have

$$P\left[\sum_{i=1}^n (u_i^2 - v) \geq a\right] \leq \exp\left\{-\frac{a^2}{2(v + ba/3)}\right\}$$

Jensen's Inequality

If g is convex on $S \subseteq \mathbb{R}$ and S is convex and closed, $P(X \in S) = 1$, $E[g(x)] < \infty$ and $EX < \infty$, then

$$E[g(X)] \geq g[E(X)].$$

Covariance inequality**Cauchy-Schwarz Inequality**

$$(EXY)^2 \leq EX^2 EY^2$$

Holder's Inequality

If $p, q > 1$ and $\frac{1}{p} + \frac{1}{q} = 1$, then

$$E|XY| \leq (E|X^p|)^{1/p} (E|X^q|)^{1/q}$$

Notation: $\|X\|_p = (E|X^p|)^{1/p}$ called the L_p norm so that

$$\|XY\| \leq \|X\|_p \|Y\|_q$$

Minokowski Inequality

For $p > 1$,

$$E(|X + Y|^p)^{1/p} \leq (E|X^p|)^{1/p} + (E|X^q|)^{1/q}$$

or

$$\|X + Y\|_p \leq \|X\|_p + \|Y\|_p$$

Lyapunov's Inequality

For any $p \geq q > 0$,

$$\|X\|_p \geq \|Y\|_q.$$

Point-wise Convergence

We say f_n converge to f point-wisely if for all x , $f_n(x) \rightarrow f(x)$.

Rmk: Not enough to ensure asymptotic property. Need stronger condition. We need to define distance d between f_n and f to be less than a very small number. If d is norm, it is norm convergence.

Norm

Let V be a collection of functions. $\|\cdot\|$ is a norm of V if

1. $\|f\| \geq 0$ for all $f \in V$
2. $\|f\| = 0$ iff f is a zero function
3. $\forall a \in \mathbb{R}, \forall f \in V, \|af\| = |a| \|f\|$
4. $\forall f, g \in V, \|f + g\| \leq \|f\| + \|g\|$

Rmk: If V satisfies these four properties, we say V is endowed with norm.

Norm Convergence

If $\|f_n - f\| \rightarrow 0$ as $n \rightarrow \infty$, we say $f_n(\cdot)$ converges to f with respect to norm $\|\cdot\|$.

Rmk: Difference definition of norm leads to different definition of convergence.

Convergence in Probability (Weak Convergence)

$\{Z_n\}$ is a sequence of random variable. Then Z_n converges to a in probability if

$$\forall \varepsilon > 0, P(\{|Z_n - a| < \varepsilon\}) \rightarrow 1 \text{ as } n \rightarrow \infty$$

or

$$\lim_{n \rightarrow \infty} P(\{|Z_n - a| < \varepsilon\}) = 1$$

and we denote this as c .

Almost Sure Convergence (Strong Convergence)

Z_n converges to a with probability 1 if

$$P\left(\left\{\lim_{n \rightarrow \infty} Z_n = a\right\}\right) = 1$$

and we denote it as $Z_n \xrightarrow{a.s.} a$.

Rmk: Almost sure event A means $P(A) = 1$.

Convergence in r -th Mean

If $E|X_n|^r < \infty$ for all n and $E|Z_n - a|^r \rightarrow 0$ as $n \rightarrow \infty$, then Z_n converges to a in r -th mean.

Rmk: When $r = 2$, we have convergence in mean square error.

Convergence in Distribution (Weakest Convergence)

$\{Z_n\}$ is a sequence of random variable with c.d.f. $\{F_n\}$ and Z is a random variable with c.d.f. F . If,

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

then sequence $\{Z_n\}$ converges in distribution to Z and we denote it as $Z_n \xrightarrow{D} Z$.

Interrationship between convergences

1. If $Z_n \xrightarrow{P} a$, then we have $Z_n \xrightarrow{D} Z$.
2. If $Z_n \xrightarrow{a.s.} a$, then we have $Z_n \xrightarrow{P} a$.
3. If Z_n converges in r -th mean, we have $Z_n \xrightarrow{P} a$.
4. If Z_n converges in k -th mean and for all $r \leq k$, we have Z_n converges in r -th mean.

Special relationship between convergences

1. If $Z_n \xrightarrow{D} Z$ where Z distribution with function which admits a constant value with probability 1, then $Z_n \xrightarrow{P} a$.
2. If $Z_n \xrightarrow{P} Z$ and $P(|Z_n| \leq k) = 1$ for all n , then Z_n converges in r -th mean.
3. If $\sum_{n=1}^{\infty} P(|Z_n - Z| > \varepsilon) < \infty$, for all $\varepsilon > 0$, then $Z_n \xrightarrow{a.s.} Z$.

Law of Large number and Central Limit Theorem

$\{X_i\}$ are i.i.d. random variables and $\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$

1. If $EX_i = \mu$ then $\bar{X} \xrightarrow{D} \mu$.
2. (WLLN) If $EX_i = \mu$ and $Var X_i = \sigma^2 < \infty$, then $\bar{X} \xrightarrow{P} \mu$.

Proof: First note that $E\bar{X}_n = \mu$ and $Var \bar{X}_n = \sigma^2/n$. By Chebyshev's Inequality, we have

$$P(|\bar{x} - \mu| \geq \varepsilon) \leq \frac{\sigma^2}{\varepsilon^2 n}$$

so that

$$1 - P(|\bar{x} - \mu| \leq \varepsilon) \leq \frac{\sigma^2}{\varepsilon^2 n}$$

or

$$P(|\bar{x} - \mu| \leq \varepsilon) \geq 1 - \frac{\sigma^2}{\varepsilon^2 n}$$

Taking $n \rightarrow \infty$, we have

$$\lim_{n \rightarrow \infty} P(|\bar{x} - \mu| \leq \varepsilon) \geq 1.$$

However, probability is bounded by 1, so we have

$$\lim_{n \rightarrow \infty} P(|\bar{x} - \mu| \leq \varepsilon) = 1.$$

3. (SLLN) If $EX_i = \mu$, then $\bar{X} \xrightarrow{a.s.} \mu$
4. (CLT) If $EX_i = \mu$ and $Var X_i = \sigma^2 < \infty$, then

$$\frac{\sqrt{n}(\bar{x} - \mu)}{\sigma} \rightarrow N(0, 1).$$

Rmk:

$$\frac{\sqrt{n}(\bar{x} - \mu)}{\sigma} = \frac{1}{\sqrt{n}} \sum_{i=1}^n \left(\frac{x_i - \mu}{\sigma} \right)$$

Chapter 2

Statistical Inference

2.1 Lecture 14

Estimation

For data set $\{X_i\}_{i=1}^n$, we wish to estimate θ , the unknown parameter of a parametric model P_θ .

Point Estimator

Given a sample $\{X_i\}_{i=1}^n$ and consider a parameter θ . Let T be a function of the sample and $\hat{\theta} = T(X_1, \dots, X_n)$ is called a point estimator.

Interval Estimator

Given $\{X_i\}_{i=1}^n$, let T_1 and T_2 be functions of sample. Let $\hat{\theta}_L = T_1(X_1, \dots, X_n)$ and $\hat{\theta}_U = T_2(X_1, \dots, X_n)$. If

$$P(\hat{\theta}_L \leq \theta \leq \hat{\theta}_U) = 1 - \alpha$$

then $[\hat{\theta}_L, \hat{\theta}_U]$ is an interval estimator with confidence probability $1 - \alpha$.

Unbiasedness

Define $bias(\hat{\theta}) = E(\hat{\theta}) - \theta$.

An unbiased estimator $\hat{\theta}$ means $bias(\hat{\theta}) = 0$ or $E(\hat{\theta}) = \theta$.

A biased estimator $\hat{\theta}$ means $bias(\hat{\theta}) \neq 0$ or $E(\hat{\theta}) \neq \theta$.

Variance of Estimator

$$Var(\hat{\theta}) = E[\hat{\theta} - E(\hat{\theta})]^2$$

Mean-square Error (MSE)

Combined measure of unbiasedness and variance of estimator.

$$\begin{aligned} MSE(\hat{\theta}) &= E(\hat{\theta} - \theta)^2 \\ &= E[\hat{\theta} - E(\hat{\theta}) + E(\hat{\theta}) - \theta]^2 \\ &= E[\hat{\theta} - E(\hat{\theta})]^2 + 2E\{[\hat{\theta} - E(\hat{\theta})][E(\hat{\theta}) - \theta]\} + E[\hat{\theta} - \theta]^2 \\ &= E[\hat{\theta} - E(\hat{\theta})]^2 + 2[E(\hat{\theta}) - \theta]E[\hat{\theta} - E(\hat{\theta})] + E[\hat{\theta} - \theta]^2 \\ &= Var(\hat{\theta}) + bias^2(\hat{\theta}) \end{aligned}$$

Consistency

If $\hat{\theta} \rightarrow \theta$ as $n \rightarrow \infty$, then $\hat{\theta}$ is a consistent estimator. In particular,

if $\hat{\theta} \xrightarrow{P} \theta$, $\hat{\theta}$ is a weak consistent estimator and

if $\hat{\theta} \xrightarrow{a.s.} \theta$, $\hat{\theta}$ is a strong consistent estimator.

Rmk: $\hat{\theta} \xrightarrow{P} \theta$ means $\forall \varepsilon > 0$, $P(|\hat{\theta} - \theta| > \varepsilon) \rightarrow 0$ as $n \rightarrow \infty$. By Markov inequality,

$$P(|\hat{\theta} - \theta| > \varepsilon) \leq \frac{E|\hat{\theta} - \theta|^2}{\varepsilon^2}.$$

Since ε is fixed, $E|\hat{\theta} - \theta|^2 = MSE(\hat{\theta}) \rightarrow 0$ as $n \rightarrow \infty$ implies consistency and so it suffices to show $bias^2(\hat{\theta}) \rightarrow 0$ and $Var(\hat{\theta}) \rightarrow 0$ as $n \rightarrow \infty$ for consistency of $\hat{\theta}$.

Sufficiency

T is a sufficient statistic if the conditional distribution of $(X_1, \dots, X_n)$ given T is independent of θ .

Rmk: Informally, it means T summarizes all information (from the sample) relevant to θ .

Factorization Theorem

Let $(X_1, \dots, X_n)$ be random sample from population whose distribution is dependent on θ . $T = T(X_1, \dots, X_n)$ is a sufficient statistic for θ if and only if the

joint density of X and θ can be written as

$$f(X_1, \dots, X_n; \theta) = h(T, \theta) H(X_1, \dots, X_n)$$

2.2 Lecture 15

Point Estimation of μ

Natural candidate: sample mean

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

Unbiasedness:

$$\begin{aligned} E(\bar{x}) &= E\left(\frac{1}{n} \sum_{i=1}^n x_i\right) \\ &= \frac{1}{n} \sum_{i=1}^n E(x_i) \\ &= \frac{1}{n} n\mu = \mu \end{aligned}$$

Hence, $\bar{x}$ is unbiased estimator.

Consistency:

$$\begin{aligned} Var(\bar{x}) &= Var\left(\frac{1}{n} \sum_{i=1}^n x_i\right) \\ &= \frac{1}{n^2} Var \sum_{i=1}^n x_i \\ &= \frac{1}{n^2} \sum_{i=1}^n Var(x_i) \\ &= \frac{1}{n^2} n\sigma^2 = \frac{\sigma^2}{n} \end{aligned}$$

therefore,

$$\begin{aligned} MSE(\bar{x}) &= bias^2(\bar{x}) + Var(\bar{x}) \\ &= 0 + \frac{\sigma^2}{n} \rightarrow 0 \text{ as } n \rightarrow \infty \end{aligned}$$

Hence, as $\bar{x} \rightarrow \mu$, $\bar{x}$ is consistent estimator.

Point Estimation of σ^2

Natural candidate: sample variance

$$s^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$$

Unbiasedness:

$$\begin{aligned} E(s^2) &= E\left[\frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2\right] \\ &= \frac{1}{n} E \sum_{i=1}^n (x_i^2 - 2x_i\bar{x} + \bar{x}^2) \\ &= \frac{1}{n} E \left[\sum_{i=1}^n x_i^2 - 2\bar{x} \sum_{i=1}^n x_i + \sum_{i=1}^n \bar{x}^2 \right] \\ &= \frac{1}{n} E \left[\sum_{i=1}^n x_i^2 - 2n\bar{x}^2 + n\bar{x}^2 \right] \\ &= \frac{1}{n} (nEx_i^2) - E\bar{x}^2 = Ex^2 - E\bar{x}^2 \end{aligned}$$

Now consider

$$\begin{aligned} E\bar{x}^2 &= E \left[\frac{\sum_{i=1}^n x_i}{n} \right]^2 \\ &= \frac{1}{n^2} E \left[\sum_{i=1}^n x_i \right]^2 \\ &= \frac{1}{n^2} E \left[\sum_{i=1}^n x_i^2 + \sum_{i \neq j} x_i x_j \right] \\ &= \frac{1}{n^2} [nEx^2 + n(n-1)ExEx] \\ &= \frac{1}{n} Ex^2 + \frac{n-1}{n} (Ex)^2 \end{aligned}$$

so that

$$\begin{aligned} E(s^2) &= Ex^2 - E\bar{x}^2 \\ &= Ex^2 - \left[\frac{1}{n} Ex^2 + \frac{n-1}{n} (Ex)^2 \right] \\ &= \frac{n-1}{n} [Ex^2 - (Ex)^2] \\ &= \frac{n-1}{n} \sigma^2 \end{aligned}$$

Hence, s^2 is biased estimator. The bias would converge to zero as n goes to infinity,

$$\begin{aligned} \text{bias}(s^2) &= E(s^2) - \sigma^2 \\ &= \frac{n-1}{n} \sigma^2 - \sigma^2 \\ &= -\frac{1}{n} \sigma^2 \\ &\rightarrow 0 \text{ as } n \rightarrow \infty \end{aligned}$$

so that s^2 is asymptotically unbiased.

Unbiased Version of sample variance would be

$$\hat{\sigma}^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

Its unbiasedness could be proved easily as

$$\begin{aligned} E(\hat{\sigma}^2) &= E\left(\frac{n}{n-1} s^2\right) \\ &= \frac{n}{n-1} E(s^2) \\ &= \frac{n}{n-1} \frac{n-1}{n} \sigma^2 \\ &= \sigma^2 \end{aligned}$$

Rmk: The reason for $n-1$ but not n can be explained by degree of freedom or using the mathematical rank concept.

Distribution of $\bar{x}$

Since $\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$, we have

$$\phi_{\hat{\mu}}(t) = \phi_{\frac{1}{n} \sum x_i}(t)$$

so that if we know the distribution of x_i , then we know the distribution of sample mean.

Even if we don't know the sample mean, by central limit theorem, we have

$$\frac{\sqrt{n}(\mu - \hat{\mu})}{\sigma} \xrightarrow{D} N(0, 1).$$

Distribution of s^2

For simplicity, we assume $X_i \stackrel{i.i.d.}{\sim} N(\mu, \sigma^2)$.

Then $\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$ would have distribution $N\left(\mu, \frac{\sigma^2}{n}\right)$.

Notice that $\frac{X_i - \mu}{\sigma} \sim N(0, 1)$ so that $\frac{(X_i - \mu)^2}{\sigma^2} \sim \chi_1^2$. We have to show $\frac{ns^2}{\sigma^2} = \sum_{i=1}^n \frac{(X_i - \mu)^2}{\sigma^2} \sim \chi_{n-1}^2$

Consider orthogonal transformation of

$$L = \begin{bmatrix} \frac{1}{\sqrt{n}} & \frac{1}{\sqrt{n}} & \frac{1}{\sqrt{n}} & \cdots & \frac{1}{\sqrt{n}} \\ \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} & 0 & \cdots & 0 \\ 0 & 0 & \ddots & & \vdots \\ \vdots & \vdots & & \ddots & \vdots \\ \frac{1}{\sqrt{n(n-1)}} & \frac{1}{\sqrt{n(n-1)}} & \cdots & \cdots & \frac{1}{\sqrt{n(n-1)}} \end{bmatrix}$$

such that $L(X - \mu) = Y$, $Y = [Y_1 \dots Y_n]^T$, $X - \mu = [X_1 - \mu \dots X_n - \mu]^T$ and $Y^T Y = [X - \mu]^T [X - \mu]$ implies

$$\sum_{i=1}^n Y_i^2 = \sum_{i=1}^n (X_i - \mu)^2$$

Jacobian matrix is

$$|J| = \left| \frac{dX}{dY} \right| = 1$$

which implies

$$\begin{aligned} Y_1 &= \sqrt{n}(\bar{X} - \mu) \\ \sum_{i=2}^n Y_i^2 &= \sum_{i=1}^n (X_i - \bar{X})^2 \end{aligned}$$

So that

$$\begin{aligned} \sum_{i=1}^n \frac{(X_i - \mu)^2}{\sigma^2} &= \sum_{i=1}^n Y_i^2 \\ &= Y_1^2 + \sum_{i=2}^n Y_i^2 \\ &= \sqrt{n}(\bar{X} - \mu)^2 + \sum_{i=1}^n (X_i - \bar{X})^2 \\ &= \chi_1^2 + \chi_{n-1}^2 \end{aligned}$$

Therefore, we have the result that

$$\frac{ns^2}{\sigma^2} = \sum_{i=1}^n \frac{(X_i - \bar{X})^2}{\sigma^2} \sim \chi_{n-1}^2$$

and s^2 is independent of $\bar{X}$ so that

$$\frac{\hat{\mu} - \mu}{\sqrt{\sigma^2/n}} \sim N(0, 1)$$

However, we usually have no information on σ^2 so we have to replace σ^2 by $\hat{\sigma}^2$. The distribution would then be

$$\begin{aligned} \frac{\hat{\mu} - \mu}{\sqrt{\hat{\sigma}^2/n}} &= \frac{(\hat{\mu} - \mu) / \sqrt{\sigma^2/n}}{\sqrt{\hat{\sigma}^2/n} / \sqrt{\sigma^2/n}} \\ &= \frac{(\hat{\mu} - \mu) / \sqrt{\sigma^2/n}}{\sqrt{\hat{\sigma}^2/\sigma^2}} \\ &= \frac{(\hat{\mu} - \mu) / \sqrt{\sigma^2/n}}{\sqrt{\frac{1}{n-1} \sum_{i=1}^n \frac{(x_i - \mu)^2}{\sigma^2}}} \\ &= \frac{N(0, 1)}{\sqrt{\chi_{n-1}^2/n - 1}} \sim t_{n-1} \end{aligned}$$

2.3 Lecture 16

Method of Moment (MOM)

Random sample $X_1, \dots, X_n$ and parametric model P_θ where $\theta = (\theta_1, \dots, \theta_k)$ is k -dimensional parameter.

Population moment

$$EX_i^r = \int x^r dF_\theta(x) = \mu_r(\theta)$$

and sample moment

$$\hat{\mu}_r(\theta) = \frac{1}{n} \sum_{i=1}^n x_i^r$$

Equate first k raw moments: For $r = 1, \dots, k$

$$\hat{\mu}_r(\theta) = \mu_r(\theta)$$

Generalized Method of Moment (GMM)

Instead of equating population moment to sample moment, we impose more general moment conditions:

$$E[h(X_i, \theta)] = 0$$

where $h_r = x_i^r - \mu_r(\theta)$ is the special case of MOM.

We have k parameter as $\theta = (\theta_1, \dots, \theta_k)$ and r restriction as $h(X_i, \theta) = [h_1(X_i, \theta), \dots, h_r(X_i, \theta)]^T$. Therefore,

if $k = r$, then the system is exactly identified;

if $k > r$, then the system is under identified;

if $k < r$, then the system is over identified.

Rmk: under-identified case is not considered usually in economics

When $k = r$, set

$$\frac{1}{n} \sum_{i=1}^n h(x_i, \theta) = 0$$

When $k < r$, we would try to minimize

$$\left\| \frac{1}{n} \sum_{i=1}^n h(x_i, \theta) \right\|$$

with a weighting matrix W , so that, we have to

$$\min_{\theta} \left[\frac{1}{n} \sum_{i=1}^n h(x_i, \theta) \right]^T W \left[\frac{1}{n} \sum_{i=1}^n h(x_i, \theta) \right]$$

Properties of GMM Estimator

Exact identification

Consistency: Yes, if it satisfies identification condition,

$$\hat{\theta} \rightarrow \theta_0, \theta \notin B_\varepsilon(\theta_0) \text{ s.t. } \frac{1}{n} \sum_{i=1}^n h(x_i, \theta) = 0$$

Distribution of $\hat{\theta}$: Assuming consistency, using Taylor's expansion,

$$\frac{1}{n} \sum_{i=1}^n h(x_i, \theta) + \frac{1}{n} \sum_{i=1}^n \frac{\partial h(x_i, \theta)}{\partial \theta} (\hat{\theta} - \theta) + O_p \left(\left\| \hat{\theta} - \theta \right\|^2 \right) = 0$$

As $\theta \rightarrow \theta$, then $O_p \left(\left\| \hat{\theta} - \theta \right\|^2 \right) \rightarrow 0$ and we have

$$\frac{1}{n} \sum_{i=1}^n h(x_i, \theta) + \frac{1}{n} \sum_{i=1}^n \frac{\partial h(x_i, \theta)}{\partial \theta} (\hat{\theta} - \theta) = 0$$

so that

$$(\hat{\theta} - \theta) \simeq \left[\frac{1}{n} \sum_{i=1}^n \frac{\partial h(x_i, \theta)}{\partial \theta} \right]^{-1} \left[\frac{1}{n} \sum_{i=1}^n h(x_i, \theta) \right]$$

If X_i are i.i.d., then by Law of large number,

$$\frac{1}{n} \sum_{i=1}^n \frac{\partial h(x_i, \theta)}{\partial \theta} \xrightarrow{P} E \left[\frac{\partial h(x_i, \theta)}{\partial \theta} \right] \equiv D$$

and

$$\frac{1}{n} \sum_{i=1}^n h(x_i, \theta) \xrightarrow{P} E[h(x_i, \theta)] = 0$$

However, given $E[h(x_i, \theta)] = 0$, then X_i are mean zero i.i.d., then by central limit theorem,

$$\frac{1}{\sqrt{n}} \sum_{i=1}^n h(x_i, \theta) \xrightarrow{D} N(0, \Sigma)$$

Hence,

$$\sqrt{n}(\hat{\theta} - \theta) \xrightarrow{D} D^{-1}N(0, \Sigma) = N(0, D^{-1}\Sigma D^{-1})$$

Over-indentification

Similar derivation with similar result.

Maximum likelihood estimation (MLE)

$(X_1, \dots, X_n)$ are i.i.d. with density function $f(x_i, \theta)$ and joint density function

$$f(x; \theta) = \prod_{i=1}^n f(x_i, \theta).$$

Consider $f(x; \theta)$ as function of θ given $(X_1, \dots, X_n)$ is called likelihood function.

$$\hat{\theta}_{MLE} = \arg \max_{\theta} L(\theta)$$

where $L(\theta) = f(x; \theta)$

Log-transformation:

$$\ell(x_i, \theta) \equiv \log L(\theta) = \sum_{i=1}^n \log f(x_i, \theta)$$

with F.O.C.

$$\frac{\partial \log L(\theta)}{\partial \theta} = 0 \Rightarrow \frac{\partial}{\partial \theta} \left[\sum_{i=1}^n \log f(x_i, \theta) \right] = 0$$

Rmk 1: GMM consider $\frac{1}{n} \sum_{i=1}^n h(x_i, \theta) = 0$

Rmk 2: score function is defined $s(x, \theta) = \frac{\partial \ell(x, \theta)}{\partial \theta}$

Rmk 3: MLE depends on whole distribution and GMM just need moment restriction

Using GMM method, we know

$$\sqrt{n}(\hat{\theta} - \theta) \rightarrow N(0, I^{-1}(\theta))$$

where

$$I(\theta) = E \left[\frac{\partial \ell(x_i, \theta)}{\partial \theta} \frac{\partial \ell(x_i, \theta)}{\partial \theta}^T \right]$$

which is called information matrix.

Properties of MLE

MLE is efficient which implies smallest variance and covariance matrix. For any other consistent estimator, say $\hat{\theta}$, if

$$\sqrt{n}(\hat{\theta} - \theta) \rightarrow N(0, V)$$

then $V \geq I^{-1}(\theta)$. Therefore, $V - I^{-1}(\theta)$ is semi-positive definite matrix.

If $\hat{\theta}$ is M.L.E., then $\hat{\beta} = g(\hat{\theta})$ is also M.L.E.

MLE for normal i.i.d random sample

Likelihood function would be

$$\ell(\mu, \sigma^2) = -\frac{n}{2} \log 2\pi - n \log \sigma - \frac{\sum (x_i - \mu)^2}{2\sigma^2}$$

with F.O.C.

$$\begin{aligned} \frac{\partial \ell(\mu, \sigma^2)}{\partial \mu} &= \frac{2 \sum (x_i - \mu)}{\sigma^2} = 0 \\ \frac{\partial \ell(\mu, \sigma^2)}{\partial \sigma^2} &= -\frac{n}{2\sigma^2} - \frac{\sum (x_i - \mu)^2}{2\sigma^4} = 0 \end{aligned}$$

so that

$$\begin{aligned} \hat{\mu}_{MLE} &= \frac{1}{n} \sum_{i=1}^n x_i = \bar{x} \\ \hat{\sigma}_{MLE}^2 &= \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 = s^2 \end{aligned}$$

2.4 Lecture 17

Confidence Interval

Let $\hat{\theta}$ be point estimate. Confidence interval $[\hat{\theta}_L, \hat{\theta}_U] \ni \theta$ with confidence $1 - \alpha$ if

$$P(\theta \in [\hat{\theta}_L, \hat{\theta}_U]) = 1 - \alpha$$

Assuming unbiasedness, we have $E(\hat{\theta}) = \theta$. Procedure to find confidence interval:

1. Find the distribution of estimator $P(\hat{\theta} < x) = F(x)$
2. Find x_L and x_U such that $F(x_L) = \alpha/2$ and $F(x_U) = \alpha/2$
3. Then we to find $P(x_L \leq \hat{\theta} \leq x_U) = 1 - \alpha$ or

$$P(\theta - x_U \leq \theta - \hat{\theta} \leq \theta - x_L) = 1 - \alpha.$$

Confidence interval for normal random sample**C.I. for $\hat{\mu}_{MLE}$**

Known σ^2 : by i.i.d. property,

$$\hat{\mu}_{MLE} = \bar{x} \sim N\left(\mu, \frac{\sigma^2}{n}\right)$$

so that

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

Hence,

$$P\left(\Phi(\alpha/2) \leq \frac{\hat{\mu} - \mu}{\sigma/\sqrt{n}} \leq \Phi^{-1}(1 - \alpha/2)\right) = 1 - \alpha$$

or

$$P\left(\hat{\mu} - \Phi^{-1}\left(1 - \frac{\alpha}{2}\right) \frac{\sigma}{\sqrt{n}} \leq \mu \leq \hat{\mu} - \Phi^{-1}\left(\frac{\alpha}{2}\right) \frac{\sigma}{\sqrt{n}}\right) = 1 - \alpha$$

We have C.I. to be

$$\left[\hat{\mu} - \Phi^{-1}\left(1 - \frac{\alpha}{2}\right) \frac{\sigma}{\sqrt{n}}, \hat{\mu} - \Phi^{-1}\left(\frac{\alpha}{2}\right) \frac{\sigma}{\sqrt{n}}\right]$$

Rmk: C.I. narrows down if σ falls or n grows.

Unknown σ^2 : by i.i.d. property

$$t_{\hat{\mu}} = \frac{\hat{\mu} - \mu}{\hat{\sigma}/\sqrt{n}} \sim t_{n-1}$$

with similar result by replace Φ^{-1} to d.f. of t distribution.

C.I. for $\hat{\sigma}^2$

Unbiased point estimate would be

$$\hat{\sigma}^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

Recall the distribution of $\hat{\sigma}^2$ would be

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

so that, taking c as inverse of d.f. of χ_{n-1}^2

$$P\left(c(\alpha/2) \leq \frac{(n-1)\hat{\sigma}^2}{\sigma^2} \leq c\left(1 - \frac{\alpha}{2}\right)\right) = 1 - \alpha$$

or

$$P\left(\frac{(n-1)\hat{\sigma}^2}{c(\alpha/2)} \leq \sigma^2 \leq \frac{(n-1)\hat{\sigma}^2}{c(1 - \alpha/2)}\right) = 1 - \alpha$$

C.I. for GMM

By C.L.T., we have

$$\sqrt{n}(\hat{\theta} - \theta) \xrightarrow{D} N(0, D^{-1}\Sigma D^{-1})$$

with standardization,

$$\sqrt{n}(D^{-1}\Sigma D^{-1})^{-1/2}(\hat{\theta} - \theta) \xrightarrow{D} N(0, I_k)$$

so that

$$\frac{\hat{\mu} - \mu}{\hat{\sigma}/\sqrt{n}} \xrightarrow{D} N(0, 1)$$

Rmk: In practice, we usually have asymptotic C.I. rather than actual C.I.

2.5 Lecture 18

Hypothesis Testing

To formulate a decision rule δ such that given the data $\{X_1, \dots, X_n\}$, it is possible to infer whether a given hypothesis is supported.

Simple Hypothesis

A hypothesis is simple if together with basic assumption, it specifies the distribution completely.

Suppose for a parametric model P_θ where $\theta \in \Theta$, if a hypothesis is $\theta = \theta_0$ so that the distribution is completely known, then such a hypothesis is simple.

Composite Hypothesis

Otherwise, it is called a composite hypothesis. With such a hypothesis, we could not know the distribution completely.

Null hypothesis

Denoted as H_0 and it is hypothesis to be tested.

Alternative hypothesis

Denoted as H_1 which is usually complement of H_0 with respect to sample space Θ .

Testing Statistics

Suppose we have a null hypothesis $H_0 : \theta \in \Theta_0$ and a testing statistics $T_n = T(X_1, \dots, X_n) \in \mathcal{A}$.

Construction of a testing procedure is to set the following rules:

1. If $T_n \in \mathcal{A}_0 \subseteq \mathcal{A}$, then reject H_0 .
2. Otherwise, if $T_n \notin \mathcal{A}_0$, then accept H_0 .

Hence, $\mathcal{A}$ is called critical region or rejection region and $\mathcal{A} - \mathcal{A}_0$ is called acceptance region.

Decision error

	Accept H_0	Reject H_0
H_0 is true	correct decision	Type I error
H_0 is not true	Type II error	correct decision

Rmk 1: Cannot simultaneously reduce both types of errors.

Rmk 2: Denote α =significance level = $P(\text{Type I error}) = P(\text{Rejection } H_0 \mid H_0)$.

Rmk 3: Denote $\beta = P(\text{Type II error}) = P(\text{Accept } H_0 \mid H_1)$.

Rmk 3: Denote $1-\beta$ =power of test = $1-P(\text{Type II error}) = 1-P(\text{Accept } H_0 \mid H_1)$.

Choosing testing procedure

1. Specifies α , which controls type I error.
2. Minimize type II error among testing statistics

$$\begin{aligned} P(\text{Type II error}) &= P(\text{Accept } H_0 \mid H_1) \\ &= 1 - P(\text{Reject } H_0 \mid H_1) \end{aligned}$$

so that $\min P(\text{Type II error})$ is equivalent to $\max P(\text{Reject } H_0 \mid H_1)$ or to maximize power of the test.

Relationship to confidence Interval for normal random sample

Known σ^2 :

Given null hypothesis $H_0 : \mu = \mu_0$, we have

$$\frac{\mu - \mu_0}{\sigma/\sqrt{n}} \sim N(0, 1)$$

so that with significance level $1 - \alpha$, we have

$$P\left(\left|\frac{\mu - \mu_0}{\sigma/\sqrt{n}}\right| \leq c\right) = 1 - \alpha$$

so that

$$P\left(\left|\frac{\mu - \mu_0}{\sigma/\sqrt{n}}\right| > c\right) = \alpha$$

which is exactly the definition of type I error if the testing procedure is to reject H_0 if

$$\left|\frac{\mu - \mu_0}{\sigma/\sqrt{n}}\right| > \Phi(\alpha)$$

Therefore, the rejection region would be

$$\left\{\left|\frac{\mu - \mu_0}{\sigma/\sqrt{n}}\right| > \Phi(\alpha)\right\}$$

Rmk: Confidence level : $1 - \text{significance level}$

Known σ^2 :

Using $\hat{\sigma}$ to replace unknown population counterpart:

$$\frac{\mu - \mu_0}{\hat{\sigma}/\sqrt{n}} \sim t_{n-1}$$

so that the rejection region would be

$$\left\{\left|\frac{\mu - \mu_0}{\sigma/\sqrt{n}}\right| > c_\alpha\right\}$$

where c_α is the inverse of distribution function of t distribution.

2.6 Lecture 19

Power function

Power function of testing procedure δ is defined as

$$\pi(\theta | \delta) = P(\text{rejecting } H_0 | \theta)$$

For simple hypotheses $H_0 : \theta = \theta_0$ and $H_1 : \theta = \theta_1$, given level of significance α ,

$$\pi(\theta_0 | \delta) = \alpha$$

and the power of test would be

$$\pi(\theta_1 | \delta) = 1 - \beta$$

Rmk: For two different testing procedures δ_1 and δ_2 , one comparison method is to compare the power of $\pi(\theta_1 | \delta_1)$ and $\pi(\theta_1 | \delta_2)$.

Neyman-Pearson Lemma

Given a sample $\{x_1, \dots, x_n\}$ and likelihood function $L(\theta)$, consider a simple hypothesis $H_0 : \theta = \theta_0$ against simple alternative $H_1 : \theta = \theta_1$. The following test is the most powerful test:

$$\text{Reject } H_0 \text{ if } \frac{L(\theta_0)}{L(\theta_1)} < c$$

Rmk 1: Since rejection region $\mathcal{A}_0$ is $\left\{ \frac{L(\theta_0)}{L(\theta_1)} < c \right\}$, the value of c could be found when α is given because

$$P\left(\left\{ \frac{L(\theta_0)}{L(\theta_1)} < c \right\} \middle| H_0\right) = \alpha$$

Rmk 2: If f is continuous, we have

$$\begin{aligned} \alpha &= P(\mathcal{A}_0 | H_0) = \int_{\mathcal{A}_0} f(x, \theta_0) dx; \\ 1 - \beta &= P(\mathcal{A}_0 | H_1) = \int_{\mathcal{A}_0} f(x, \theta_1) dx \end{aligned}$$

Rmk 3: The meaning of most powerful test means that if we have another testing procedure with rejection region $\mathcal{B}_0$, we would have (i) $P(\mathcal{B}_0 | H_0) = \alpha$ and (ii) $P(\mathcal{B}_0 | H_1) \leq P(\mathcal{A}_0 | H_1)$.

Significance level of composite null hypothesis

Suppose $H_0 : \theta \in \Theta_0$ and $H_1 : \theta = \theta_1$, if for all $\theta \in \Theta_0$, such that

$$P_\theta(\text{Type I error} | H_0) = \int_{\mathcal{A}_0} f(x, \theta) dx = \pi(\theta) < \alpha$$

then, it has significance level of α .

Power of composite alternative hypothesis

Suppose $H_0 : \theta = \theta_0$ and $H_1 : \theta \in \Theta_1$, if for all $\theta \in \Theta_1$, such that

$$P_\theta(\mathcal{A}_0 | \theta) = \int_{\mathcal{A}_0} f(x, \theta) dx = \pi(\theta) = 1 - \beta$$

then, it has power of $1 - \beta$.

Uniformly most powerful Test (UMP Test)

Suppose $H_0 : \theta \in \Theta_0$ and $H_1 : \theta \in \Theta$, testing procedures δ_i , test δ^* is UMP test if for all level of α such that for the

1. same level of significance α ,

$$\pi(\theta \mid \delta_i) \leq \alpha, \text{ for all } \theta \in \Theta_0$$

2. δ^* is highest power among δ_i ,

$$\pi(\theta \mid \delta^*) \geq \pi(\theta \mid \delta_i), \text{ for all } \theta \in \Theta_1$$

Monotone Likelihood Ratio

Let $f_n(x \mid \theta)$ be joint p.d.f. of $\{x_1, \dots, x_n\}$ and consider a statistic $T = T(x_1, \dots, x_n)$ if for two values θ_1 and θ_2 in Θ , say $\theta_1 < \theta_2$, the likelihood ratio

$$\frac{f_n(x \mid \theta_2)}{f_n(x \mid \theta_1)}$$

depends on x only through $T(x)$, and this ratio is an increasing function of $T(x)$ over the range of all possible values of $T(x)$, then $f_n(x \mid \theta)$ has a monotone likelihood ratio in statistics T .

Testing with monotone likelihood ratio

For $H_0 : \theta \leq \theta_0$ and $H_1 : \theta > \theta_0$:

Suppose $f_n(x \mid \theta)$ has a monotone likelihood ratio in $T(x)$. Let c and α be constant such that $P(T \geq c \mid \theta = \theta_0) = \alpha$. Then the test “rejects $H_0 : \theta \leq \theta_0$ if $T \geq c$ ” is a UMP test at the level of significance α .

For $H_0 : \theta \geq \theta_0$ and $H_1 : \theta < \theta_0$:

Suppose $f_n(x \mid \theta)$ has a monotone likelihood ratio in $T(x)$. Let c and α be constant such that $P(T \leq c \mid \theta = \theta_0) = \alpha$. Then the test “rejects $H_0 : \theta \geq \theta_0$ if $T \leq c$ ” is a UMP test at the level of significance α .

Appendix A

The Appendix

Now the Appendix is empty.

Afterword

The back matter often includes one or more of an index, an afterword, acknowledgements, a bibliography, a colophon, or any other similar item. In the back matter, chapters do not produce a chapter number, but they are entered in the table of contents. If you are not using anything in the back matter, you can delete the back matter TeX field and everything that follows it.

Acknowledgement

Under construction

Bibliography

In writing the notes, I have consulted the classnotes by Prof. Xiao and textbook “Probability and Statistics” by DeGroot.