

MATH 357:

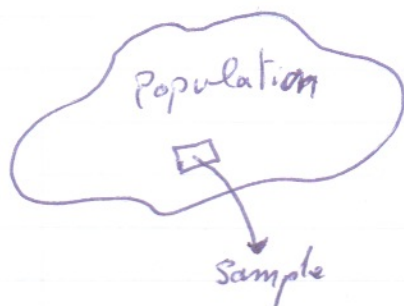
Hon. Statistics

Math 357: Hon. Statistics

Introduction:

The main goal of Statistics:

make inference (an educated guess) about a population based on a sample drawn from that population.



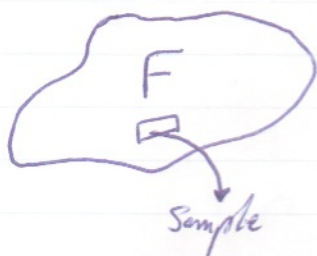
Examples:

- Say something about how long people with Alzheimer disease live from onset ~~of~~ their disease, based on a sample of people with Alzheimer.
- Say something about the proportion of Canadians that are in favor of capital punishment, based on a sample of 1000 Canadians.
- Say something about the reliability of certain aircraft components based on tests carried out on a sample of 100 such components.

We need to be more precise by what we mean by a population

First, we shall mean a population of real numbers (or vectors)

Second, when we talk of describing a population, we shall mean describing the probability distribution F of that population (of these possible values — real nb / vectors)



There are 2 main branches of Statistics:

(1) $\rightarrow$ Parametric inference

(2) $\rightarrow$ Non-parametric / distribution free inference

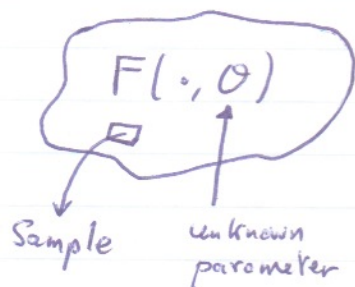
$\rightarrow$ Under (1), we begin by assuming that F belongs to some so-called parametric family of distributions. (e.g. F comes from a normal dist. or Exp dist, or Poisson)

What is unknown is one or more parameters that tie down the specific distribution in the family (e.g. $F \sim N(\mu, \sigma^2)$ where μ is unknown $\rightarrow$ inference about μ

or $F \sim N(\mu, \sigma^2)$ where the vector (μ, σ^2) is unknown $\rightarrow$ inference about $\Theta := (\mu, \sigma^2)$

or $F \sim \text{Bernoulli}(p)$ where p is unknown $\rightarrow$ inference about p)

$\rightarrow$ Under (2), we do not start out by making any assumption about the distribution F , i.e. it is not assumed to belong to any particular family.



Our inference will be entirely parametric.

Order Statistics

I Definition & Main theorem:

Order Statistics = let $X_1, \dots, X_n$ be IID r.v. drawn from some distribution F :

- the order statistics of $(X_k)_{k=1}^n$ is the sample $Y_1 \leq Y_2 \leq \dots \leq Y_n$,
- $Y_i = i^{\text{th}}$ order statistic $\leadsto \forall \omega \in \Omega: Y_i(\omega) = \min(\{X_1(\omega), \dots, X_n(\omega)\})$

$$\begin{cases} Y_1 = \min(X_1, \dots, X_n) \\ Y_2 = 2^{\text{nd}} \min(X_1, \dots, X_n) \\ \vdots \\ Y_n = \max(X_1, \dots, X_n) \end{cases}$$

Note = If the X_i 's are (absolutely) continuous with PDF f_X , then with probability 1: $Y_1 < Y_2 < \dots < Y_n$
 (8) $Y_i = Y_{i+1} \Leftrightarrow X_i = X_{i+1} \rightarrow$ happens with prob 0 ($\because X_i$'s are IID r.v.)

Main theorem: $X_1, \dots, X_n$ continuous IID r.v. with common PDF f_X :

$$\text{then: } f_{Y_1, \dots, Y_n}(y_1, \dots, y_n) = \begin{cases} n! \prod_{i=1}^n f_X(y_i), & \text{for } y_1 < y_2 < \dots < y_n \\ 0, & \text{elsewhere} \end{cases}$$

Note = Δ the Y_i 's are clearly NOT IID.

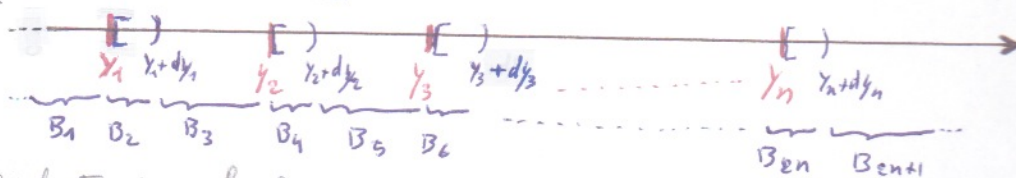
II Heuristic derivation: Multinomial distribution setup

Multinomial distribution: n independent trials, each of which can result in one of K possible outcomes.

- P_i = prob (outcome i at a trial) $\rightarrow$ constant for every trial
- X_i = # of trials resulting in outcome i

$$\Rightarrow \begin{cases} P(X_1 = n_1, \dots, X_K = n_K) = \frac{n!}{n_1! \dots n_K!} P_1^{n_1} \dots P_K^{n_K} \\ \sum_{i=1}^K n_i = n \quad \& \quad \sum_{i=1}^K P_i = 1 \end{cases}$$

Heuristic Proof =



Divide the horizontal axis into $n+1$ bins

Note: each X_i can only fall in one of these bins.

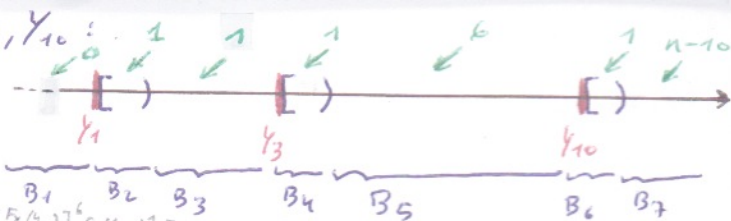
- X_i 's independent $\rightarrow$ outcomes of X_i are indep.
- X_i 's identically dist $\Rightarrow$ prob of falling in bin j is constant $\forall j$

$\Rightarrow$ Multinomial Setup!

$$\begin{aligned} \text{Joint PDF: } f(y_1, \dots, y_n) &= \lim_{dy_1, \dots, dy_n \rightarrow 0} \left[\frac{P(Y_1 \in [y_1, y_1 + dy_1], \dots, Y_n \in [y_n, y_n + dy_n])}{dy_1 \cdot dy_2 \cdot \dots \cdot dy_n} \right] \\ &= \lim_{dy_1, \dots, dy_n \rightarrow 0} \left[\frac{P(0 X_i \text{ in } B_1, 1 X_i \text{ in } B_2, 0 X_i \text{ in } B_3, \dots, 1 X_i \text{ in } B_n, 0 X_i \text{ in } B_{n+1})}{dy_1 \cdot dy_2 \cdot \dots \cdot dy_n} \right] \\ &= \lim_{dy_1, \dots, dy_n \rightarrow 0} \left[\frac{n!}{0! 1! 0! 1! \dots 1! 0!} \cdot \frac{f_X(y_1) \cdot (f_X(y_2) dy_2)^1 \cdot (f_X(y_2) - f_X(y_2))^0 \cdot \dots \cdot (f_X(y_n) dy_n)^1 \cdot (1 - f_X(y_n))^0}{dy_1 \cdot dy_2 \cdot \dots \cdot dy_n} \right] \\ &= n! \prod_{i=1}^n f_X(y_i) \end{aligned}$$

$$\Rightarrow f(y_1, \dots, y_n) = \begin{cases} n! \prod_{i=1}^n f_X(y_i), & \text{for } y_1 < y_2 < \dots < y_n \\ 0, & \text{otherwise} \end{cases}$$

Ex: Given $X_1, \dots, X_n$, find the joint PDF of Y_1, Y_3, Y_{10} :



$$f(y_1, y_3, y_{10}) = \frac{P(Y_1 \leq y_1, Y_3 \leq y_3, Y_{10} \leq y_{10})}{dy_1 dy_3 dy_{10}}$$

$$= \frac{n!}{0!1!1!6!1!(n-10)!} F_X(y_1) \cdot f_X(y_1) \cdot [F_X(y_3) - F_X(y_1)] \cdot f_X(y_3) \cdot [F_X(y_{10}) - F_X(y_3)] \cdot f_X(y_{10}) \cdot [1 - F_X(y_{10})]^{n-10}$$

$$\Rightarrow f(y_1, y_3, y_{10}) = \begin{cases} \frac{n!}{6!(n-10)!} f_X(y_1) f_X(y_3) f_X(y_{10}) \cdot [F_X(y_3) - F_X(y_1)] \cdot [F_X(y_{10}) - F_X(y_3)] \cdot [1 - F_X(y_{10})]^{n-10}, & \text{if } y_1 < y_3 < y_{10} \\ 0, & \text{otherwise} \end{cases}$$

Rigorous derivation: n -dim transformation from the X_i 's to Y_i 's

Idea: Let $\tilde{X} = (X_1, \dots, X_n)$, $\tilde{Y} = (Y_1, \dots, Y_n)$:

make a n -dim transfo $\tilde{g}: \tilde{X} \rightarrow \tilde{Y}$, $\tilde{g}(\tilde{x}) = \tilde{y}$ where $g_i(\tilde{x}) = y_i$

Case 1: If $\tilde{g}$ is 1-1: $f_{\tilde{Y}}(\tilde{y}) = f_{\tilde{X}}(g_1^{-1}(\tilde{y}), \dots, g_n^{-1}(\tilde{y})) \cdot |J|$, $J = \det\left(\frac{\partial x_i}{\partial y_j}\right)$

Case 2: If $\tilde{g}$ is not 1-1 on the support R of $f_{\tilde{X}}$, but $\tilde{g}$ is 1-1 on subregions R_e of R

$$\left. \begin{aligned} \tilde{g}_e &:= \text{restriction of } \tilde{g} \text{ to } R_e \\ J_e &:= \det\left(\frac{\partial x_i}{\partial y_j} \Big|_{x_i \in R_e}\right) \end{aligned} \right\} \Rightarrow f_{\tilde{Y}}(\tilde{y}) = \sum_{e=1}^L f_{\tilde{X}}(\tilde{g}_e(\tilde{y})) \cdot |J_e|$$

$\begin{cases} R = \bigcup_{e=1}^L R_e \\ R_e \cap R_{e'} = \emptyset \text{ for } e \neq e' \end{cases}$

Return to proof:

• the transfo $\begin{cases} y_1 = \min(x_1, \dots, x_n) \\ \vdots \\ y_n = \max(x_1, \dots, x_n) \end{cases}$ is NOT 1-1 ($\because$) Permuting the x_i 's leave the y_j 's unchanged.

• Define the region $R_e :=$ the particular permutation of x_i 's such that: $\begin{cases} y_1 = x_{e,1} \\ \vdots \\ y_n = x_{e,n} \end{cases} \Rightarrow$ the transfo restricted to R_e is 1-1!

For any other permutation: $|J_e| = 1$ ($\because$) all rows have one 1 + all the rest is 0

$\rightarrow L = \#$ of possible regions = $\#$ of possible permutations of the x_i 's $\Rightarrow L = n!$

• So: $f_{\tilde{Y}}(\tilde{y}) = \sum_{e=1}^{n!} f_{\tilde{X}}(y_1, \dots, y_n) \cdot 1 = \sum_{\substack{e=1 \\ \text{I.I.D}}}^{n!} \prod_{i=1}^n f_X(y_i) \Rightarrow f_{\tilde{Y}}(\tilde{y}) = \begin{cases} n! \prod_{i=1}^n f_X(y_i), & y_1 < y_2 < \dots < y_n \\ 0, & \text{otherwise} \end{cases}$

IV Application:

Q: $\{X_i\}_{i=1}^4$ are IID $\exp(\beta)$ r.v. with PDF: $f_X(x) = \begin{cases} \frac{1}{\beta} e^{-\frac{x}{\beta}}, & x \geq 0 \\ 0, & x < 0 \end{cases}$

Find the joint PDF of $Y_1 \times Y_4$:

A1: Multinomial way:

$$f_{Y_1, Y_4}(y_1, y_4) = \frac{4!}{0!1!2!1!0!} (F_X(y_1))^0 \cdot (f_X(y_1))^1 \cdot [F_X(y_4) - F_X(y_1)]^2 \cdot (f_X(y_4))^1 \cdot [1 - F_X(y_4)]^0$$

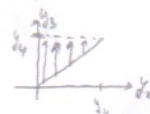
$$\Rightarrow f_{Y_1, Y_4}(y_1, y_4) = \begin{cases} \frac{2}{\beta^2} e^{-\frac{y_1+y_4}{\beta}} \cdot [e^{-\frac{y_1}{\beta}} - e^{-\frac{y_4}{\beta}}]^2, & y_1 < y_4 \\ 0, & \text{elsewhere} \end{cases}$$

A2: Integrate over the others:

$$f_{Y_1, Y_4}(y_1, y_4) = \begin{cases} 4! \int_{y_1}^{y_4} \int_{y_2}^{y_3} \frac{e^{-y_1/\beta}}{\beta} \cdot \frac{e^{-y_2/\beta}}{\beta} \cdot \frac{e^{-y_3/\beta}}{\beta} \cdot \frac{e^{-y_4/\beta}}{\beta}, & y_1 < y_4 \\ 0, & \text{elsewhere} \end{cases}$$

($\because$) $y_1 < y_2 < y_3 < y_4$

$\Rightarrow \begin{cases} y_1 < y_2 < y_4 \\ y_2 < y_3 < y_4 \end{cases}$



$\hookrightarrow$ Same result as above, but painful!

Sampling distributions :

① Sample Mean & Sample Variance :

Sample Mean = $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$
Sample Variance = $S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$

where X_i 's are IID

Sampling distribution: The distribution of functions of the observations $X_1, \dots, X_n$.

Prop: $X_1, \dots, X_n$ IID $N(\mu, \sigma^2)$ r.v.: $\frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \sim N(0, 1)$ (use MGF!)
 Note: exact relation, not a limit!

Prop: $X_1, \dots, X_n$ IID r.v. with mean μ & variance σ^2 :

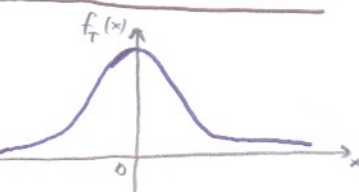
$\frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \xrightarrow{D} Z \sim N(0, 1)$ as $n \rightarrow \infty$ (use CLT!) (Convergence in distribution: the CDF converges to CDF of limit F_Z at every continuity point x of F_Z)
 $\frac{\bar{X} - \mu}{S/\sqrt{n}} \xrightarrow{D} Z \sim N(0, 1)$ as $n \rightarrow \infty$ (Slutsky: $\frac{X_n}{Y_n} \xrightarrow{D} \frac{a}{b}$ if $X_n \xrightarrow{D} a$ and $Y_n \xrightarrow{P} b \neq 0$)

Prop: $X_1, \dots, X_n$ IID $\text{Ber}(p)$ r.v.: $X = \sum_{i=1}^n X_i$, $\hat{p} = \frac{X}{n}$ Note: $X = \# \text{ of Success} \Rightarrow$ Binomial distrib°

$\frac{\frac{X}{n} - p}{\sqrt{p(1-p)/n}} \xrightarrow{D} Z \sim N(0, 1)$ as $n \rightarrow \infty$ (result 2. a with $\mu=p, \sigma^2=p(1-p)$)
 $\Rightarrow$ Normal approximation to Binomial distrib°
 $\frac{\frac{X}{n} - p}{\sqrt{p(1-p)/n}} \xrightarrow{D} Z \sim N(0, 1)$ as $n \rightarrow \infty$ (Slutsky: $\hat{p} = \frac{X}{n} \xrightarrow{P} p \Rightarrow \begin{cases} Y_n = \frac{\hat{p}(1-\hat{p})}{p(1-p)} \xrightarrow{P} 1 \\ Z_n = \frac{\frac{X}{n} - p}{\sqrt{p(1-p)/n}} \xrightarrow{D} N(0, 1) \end{cases}$)

Prop: $X_1, \dots, X_n$ IID $N(\mu, \sigma^2)$ r.v.: $\frac{\bar{X} - \mu}{S/\sqrt{n}} \sim t_\nu \Rightarrow$ (Student) t distribution with ν degrees of freedom
 Proof = Next Section!

Student t distribution: $f_{T_\nu}(x) = \frac{\Gamma(\frac{1+\nu}{2})}{\sqrt{\pi\nu} \Gamma(\frac{\nu}{2})} \cdot (1 + \frac{x^2}{\nu})^{-\frac{1+\nu}{2}} \quad \forall x \in \mathbb{R}$

$f_T(x)$ 

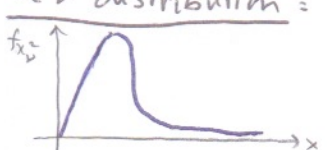
- Symmetric about zero $\forall \nu$
- $Y \sim t_1$ has a Cauchy distribution
- $Y \sim t_\nu$ does not have a ν^{th} moment: Y has moments of order $r < \nu \Rightarrow E(|Y|^r) < \infty$

$\rightarrow f_Y(y)$ has a heavier tail than the normal distrib° (the normal dist has all its moments)

χ^2_ν distribution: $f_{\chi^2_\nu}(x) = \frac{1}{2^{\frac{\nu}{2}} \Gamma(\frac{\nu}{2})} x^{\frac{\nu}{2}-1} e^{-\frac{x}{2}} \quad \forall x \geq 0 \Rightarrow$ Gamma distrib°: $\begin{cases} \alpha = \frac{\nu}{2} \\ \beta = 2 \end{cases}$

$M_{\chi^2_\nu}(t) = \frac{1}{(1-2t)^{\nu/2}}, \quad \forall t < \frac{1}{2} \Rightarrow M_{\text{Gamma}(\alpha, \beta)}(t) = \frac{1}{(1-\beta t)^\alpha} \quad \forall t < \frac{1}{\beta}$

$f_{\chi^2_\nu}(x)$ is NOT symmetric about zero!

$f_{\chi^2_\nu}(x)$ 

Prop: $X_1, \dots, X_n$ independent $\chi^2_{\nu_i}$ r.v.: $\sum_{i=1}^n X_i \sim \chi^2_{\sum \nu_i}$ (use MGF!)

Student T distribution:

Theorem = $X_1, \dots, X_n$ IID $N(\mu, \sigma^2)$ r.v.:

$$T = \frac{\bar{X} - \mu}{S/\sqrt{n}} \sim t_{n-1}$$

Lemma 1 = Let $\begin{cases} Y \sim \chi^2_v \\ Z \sim N(0,1) \end{cases}$: then if $Y \perp Z$:

$$\frac{Z}{\sqrt{Y/v}} \sim t_v$$

$$\frac{N(0,1)}{\sqrt{\chi^2_v/v}} \sim t_v$$

Proof = See HW 1

Lemma 2 = $X_1, \dots, X_n$ IID $N(\mu, \sigma^2)$ r.v. : then $\frac{(n-1)S^2}{\sigma^2} \sim \chi^2_{n-1}$ when $\bar{X} \perp S^2$

$$\text{Proof} = \frac{1}{\sigma^2} \sum_{i=1}^n (X_i - \mu)^2 = \frac{1}{\sigma^2} \sum_{i=1}^n (X_i - \bar{X} + \bar{X} - \mu)^2 = \frac{1}{\sigma^2} \sum_{i=1}^n (X_i - \bar{X})^2 + \frac{(\bar{X} - \mu)^2}{\sigma^2/n}$$

Recall:
 $Z \sim N(0,1)$
 $\Rightarrow Z^2 \sim \chi^2_1$

$$\text{So } \sum_{i=1}^n \underbrace{\left(\frac{X_i - \mu}{\sigma}\right)^2}_{\sim \chi^2_1 \text{ (as } X_i \text{'s indep)}} \sim \chi^2_n \Rightarrow \underbrace{\sum_{i=1}^n \left(\frac{X_i - \mu}{\sigma}\right)^2}_{W \sim \chi^2_n} = \underbrace{\frac{(n-1)S^2}{\sigma^2}}_{Y \sim ?} + \underbrace{\left(\frac{\bar{X} - \mu}{\sigma/\sqrt{n}}\right)^2}_{U \sim \chi^2_1}$$

• So $W = Y + U$: Lemma 3 $\Rightarrow \bar{X} \perp S^2 \Rightarrow Y \perp U$

$$\text{So } M_W(t) = M_{Y+U}(t) \stackrel{Y \perp U}{=} M_Y(t) M_U(t) \Rightarrow \frac{1}{(1-2t)^{n/2}} = M_Y(t) \cdot \frac{1}{(1-2t)^{1/2}} \Rightarrow M_Y(t) = \frac{1}{(1-2t)^{(n-1)/2}}$$

Lemma 3 = $X_1, \dots, X_n$ IID $N(\mu, \sigma^2)$ r.v. : then $\bar{X} \perp S^2 \Rightarrow Y = \frac{(n-1)S^2}{\sigma^2} \sim \chi^2_{n-1}$

Proof = Idea: Write S^2 as a function of $(X_1 - \bar{X}, \dots, X_n - \bar{X}) \Rightarrow$ show: $\bar{X} \perp (X_1 - \bar{X}, \dots, X_n - \bar{X})$ (only true for Normal dist)

WLOG: $\mu=0, \sigma^2=1$ (reScale the Normal dist) $\sum_{i=1}^n (X_i - \bar{X}) = 0$

$$\bullet S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{1}{n-1} \left((X_1 - \bar{X})^2 + \sum_{i=2}^n (X_i - \bar{X})^2 \right) = \frac{1}{n-1} \left[\sum_{i=2}^n (X_i - \bar{X}) \right] + \frac{1}{n-1} \sum_{i=2}^n (X_i - \bar{X})^2$$

$$\bullet \text{transfo: } \begin{cases} Y_1 = \bar{X} \\ Y_i = X_i - \bar{X} \quad \forall 2 \leq i \leq n \end{cases} \Rightarrow \begin{cases} Y_1 = \bar{X} \\ Y_i = X_i - \bar{X} \quad \forall 2 \leq i \leq n \end{cases} \Rightarrow \begin{cases} X_1 = nY_1 - \sum_{i=2}^n (Y_1 + Y_i) \\ X_i = Y_1 + Y_i \quad \forall 2 \leq i \leq n \end{cases}$$

$$\Rightarrow \begin{cases} X_1 = Y_1 - \sum_{i=2}^n Y_i \\ X_i = Y_1 + Y_i \quad \forall 2 \leq i \leq n \end{cases} \text{ is a 1-1 Transfo} \Rightarrow |J| = \left| \det \begin{pmatrix} 1 & -1 & -1 & \dots & -1 \\ 1 & 1 & 0 & \dots & 0 \\ 1 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & 0 & 0 & \dots & 1 \end{pmatrix} \right| = \left| \det \begin{pmatrix} n & 0 & 0 & \dots & 0 \\ 1 & 1 & 0 & \dots & 0 \\ 1 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & 0 & 0 & \dots & 1 \end{pmatrix} \right|$$

$$\bullet \text{So: } f_{Y_1, \dots, Y_n}(y_1, \dots, y_n) = f_{X_1, \dots, X_n}(x_1, \dots, x_n) \cdot |J|$$

$$[X_i \text{ are IID}] = \left(\prod_{i=1}^n f_X(x_i) \right) \cdot n = \frac{n}{(2\pi)^{n/2}} e^{-\frac{1}{2} \sum_{i=1}^n x_i^2} = \frac{n}{(2\pi)^{n/2}} e^{-\frac{1}{2} (x_1^2 + \sum_{i=2}^n x_i^2)} = c e^{-\frac{1}{2} (y_1 - \sum_{i=2}^n y_i)^2 + \sum_{i=2}^n (y_1 + y_i)^2}$$

$$= c e^{-\frac{1}{2} n y_1^2} e^{-\frac{1}{2} \left[\left(\sum_{i=2}^n y_i \right)^2 + \sum_{i=2}^n y_i^2 \right]}$$

$$\Rightarrow f_{Y_1, \dots, Y_n}(y_1, \dots, y_n) = c g_1(y_1) g_2(y_2, \dots, y_n) \Rightarrow Y_1 \perp (Y_2, \dots, Y_n) \Rightarrow \bar{X} \perp (X_1 - \bar{X}, \dots, X_n - \bar{X}) \Rightarrow \bar{X} \perp S^2$$

$$\text{Proof} = T = \frac{\bar{X} - \mu}{S/\sqrt{n}} = \frac{(\bar{X} - \mu)/(\sigma/\sqrt{n})}{\sqrt{\frac{(n-1)S^2}{\sigma^2}/(n-1)}} \sim \frac{Z}{\sqrt{Y/v}} \sim t_v \text{ where: } \begin{cases} Z \sim N(0,1) & (\because \bar{X} \sim N(\mu, \sigma^2)) \\ Y \sim \chi^2_{n-1} & (\because \text{Lemma 2 + 3}) \end{cases}$$

III F-distribution :

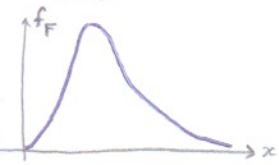
Thm : $(\bar{X}_{n_1}, s_{n_1}^2) \perp (\bar{Y}_{n_2}, s_{n_2}^2)$ r.v. with $\begin{cases} X_i \text{'s IID } N(\mu_1, \sigma^2) \\ Y_i \text{'s IID } N(\mu_2, \sigma^2) \end{cases}$ Note: Same σ^2 !

Then : $\frac{S_1^2}{S_2^2} \sim F_{n_1-1, n_2-1}$ with : $\begin{cases} S_1^2 = \frac{1}{n_1-1} \sum_{i=1}^{n_1} (X_i - \bar{X})^2 \\ S_2^2 = \frac{1}{n_2-1} \sum_{i=1}^{n_2} (Y_i - \bar{Y})^2 \end{cases}$

Proof : $\frac{S_1^2}{S_2^2} = \frac{(n_1-1)S_1^2}{(n_1-1)\sigma^2} / \frac{(n_2-1)S_2^2}{(n_2-1)\sigma^2} \stackrel{\text{indep!}}{\sim} \frac{\chi_{n_1-1}^2/n_{1-1}}{\chi_{n_2-1}^2/n_{2-1}} \underset{\text{Lemma}}{\sim} F_{n_1-1, n_2-1}$ (o.o) $\bar{X}_1 \perp Y_1 \Rightarrow S_1^2 \perp S_2^2 \Rightarrow \frac{(n_1-1)S_1^2}{\sigma^2} \perp \frac{(n_2-1)S_2^2}{\sigma^2}$

F-distribution = F_{ν_1, ν_2} : F distribution with degrees of freedom ν_1 & ν_2 .

Prop :



$$f_{F_{\nu_1, \nu_2}}(x) = \begin{cases} \frac{\Gamma(\frac{\nu_1 + \nu_2}{2})}{\Gamma(\frac{\nu_1}{2})\Gamma(\frac{\nu_2}{2})} \left(\frac{\nu_1}{\nu_2}\right)^{\frac{\nu_1}{2}} \cdot \frac{x^{\frac{\nu_1}{2}-1}}{(1 + \frac{\nu_1}{\nu_2}x)^{\frac{\nu_1 + \nu_2}{2}}} & , 0 \leq x < \infty \\ 0 & , x < 0 \end{cases}$$

- $Y \sim F_{\nu_1, \nu_2}$ has a support over $[0, \infty)$
- PDF of F_{ν_1, ν_2} is skewed to the right

Lemma = If $\begin{cases} U \sim \chi_{\nu_1}^2 \\ V \sim \chi_{\nu_2}^2 \end{cases}$ and $U \perp V$, then

$$\boxed{\frac{U/\nu_1}{V/\nu_2} \sim F_{\nu_1, \nu_2}}$$

$$\frac{\chi_{\nu_1}^2/\nu_1}{\chi_{\nu_2}^2/\nu_2} \sim F_{\nu_1, \nu_2}$$

Proof = HW 1

STATISTICAL

INFERENCE:

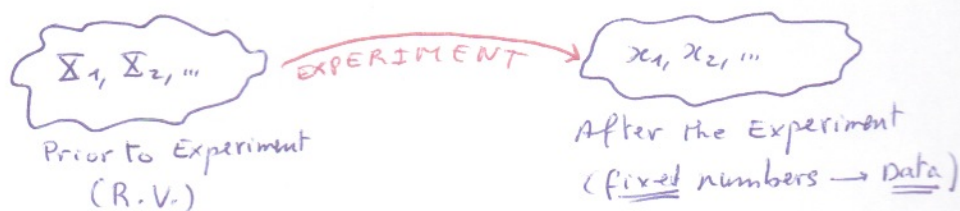
Assume that we have observations $X_1, \dots, X_n$ drawn from some distribution $F(\cdot; \theta)$, where θ is an unknown parameter:

2 main problems:

- ① Estimate θ : make an educated guess about θ .
- ② Hypothesis testing: carry out a test of hypothesis about θ .

Terminology:

Idea:



- Prior to the "experiment" (i.e. before observing our data):
the quantities $X_1, X_2, \dots$ are R.V.'s $\rightarrow$ we don't know yet what values they will assume!
- After the "experiment":
the realized values of $X_1, X_2, \dots$ are denoted by $x_1, x_2, \dots \rightarrow$ Fixed numbers!

Statistic = any measurable function $T(X_1, X_2, \dots, X_n)$ Ex: $\bar{X}, s^2, X_{26}, 3.62, \frac{\bar{X}-\mu}{s/\sqrt{n}}, \frac{\bar{X}-\mu}{s/\sqrt{n}}, \dots$

Estimator = $\hat{\theta} = \hat{\theta}(X_1, \dots, X_n)$ is an estimator for θ if $\hat{\theta}$ is a statistic that does NOT depend functionally on any unknown parameters
 $\rightarrow \hat{\theta}$ is used to estimate θ !

- Note:
- Before the experiment, $\hat{\theta}(X_1, \dots, X_n)$ is a R.V. ($\because \hat{\theta}$ = function of R.V.'s)
 $\Rightarrow \hat{\theta}$ has a probability distribution: call it its "(sampling) distribution".
 - After the experiment, $\hat{\theta}(x_1, \dots, x_n)$ is a constant (number or vector).
 $\Rightarrow \hat{\theta}(x_1, \dots, x_n)$ is called an ESTIMATE of θ .

Ex = Given $X_1, \dots, X_n$:

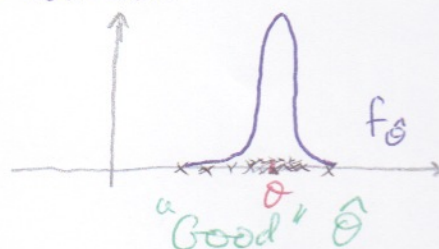
- $\hat{\theta} = X_1$ is an estimator for $\mu, \sigma^2, \dots$
- $\hat{\theta} = X_1 + 3\mu$ is NOT estimator for μ ($\because \hat{\theta}$ depends on μ)
- $\hat{\theta} = \frac{X_1 + 2}{\sigma^2}$ is an estimator for μ only if we know σ^2 .

"Good" estimators = we seek recipes that work well on average when we use it!
However, once the experiment is done, we don't know how well we have done:
we can only appeal to the fact that we used a recipe that works well on average.

Idea = Desirable to have a (sampling) distribution $f_{\hat{\theta}}$ of $\hat{\theta}$ closely concentrated about the parameter θ that we want to estimate.



$\rightarrow$ Most of the time, the $\hat{\theta}$ we will get will be far from θ because the sampling dist not concentrated about θ



$\hookrightarrow$ Sampling dist concentrated about θ : most of the time, the $\hat{\theta}$ we get will be close to θ .

the summary page =
the procedures used are only
"good" in the sense that they are
"good on average" if the procedure
was used repeatedly: we can never say how well we have done on a particular occasion!

Unbiasedness :

→ one of the criteria of "goodness"
that estimators should share!

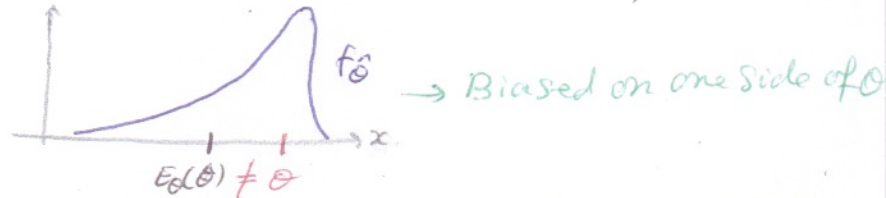
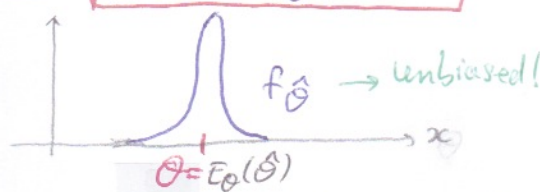
I Definition :

Unbiased Estimator = Let $X_1, \dots, X_n$ be RV's drawn from some distribution with parameter θ

$\hat{\theta}$ is an unbiased estimator for θ if $E_{\theta}(\hat{\theta}) = \theta, \forall \theta \in \Theta$ (Note: Subscript of E_{θ} : use the distribⁿ of $(X_1, \dots, X_n)$ at the parameter θ when calculating the expected value)

Parameter Space Θ : Set of all possible values of θ

Bias : $b(\theta) := E_{\theta}(\hat{\theta}) - \theta$



Note: usually start with a $\hat{\theta}$ proposal & check if biased: it is often possible to correct for the bias by redefining $\hat{\theta}$.

II Examples :

Ex 1 : $X_1, \dots, X_n$ IID R.V. : let the mean μ be unknown

Prop : $\hat{\mu} = \bar{X}$ is an unbiased estimator for μ (or) $E_{\mu}(\hat{\mu}) = E_{\mu}(\bar{X}) = \mu \forall \mu$ by def of μ

Ex : $\hat{\mu} = X_1$ (or) $E_{\mu}(X_1) = \mu$ by def, $\hat{\mu} = \frac{X_1 + X_2}{2}$ (or) $E_{\mu}(\hat{\mu}) = \frac{E(X_1) + E(X_2)}{2} = \frac{\mu + \mu}{2} = \mu$

Bias : $\hat{\mu} = \bar{X} + 3$ (Correct: $\hat{\mu} \rightarrow \hat{\mu} - 3$) (or) $E_{\mu}(\bar{X} + 3) = \mu + 3 \neq \mu \forall \mu$
 $\hat{\mu} = 3$ (or) $E_{\mu}(\hat{\mu}) = E_{\mu}(3) = 3 \neq \mu \forall \mu \neq 3$

Ex 2 : $X_1, \dots, X_n$ IID R.V. with finite mean μ and variance σ^2 :

Prop : μ known, σ^2 unknown : $\hat{\sigma}^2 = \frac{\sum_{i=1}^n (X_i - \mu)^2}{n}$ is unbiased for σ^2

$$\begin{aligned} \text{(or) } E_{\sigma^2}(\hat{\sigma}^2) &= \frac{1}{n} \sum_{i=1}^n E_{\sigma^2}((X_i - \mu)^2) \\ &= \frac{1}{n} \sum_{i=1}^n (\sigma^2 + E(X_i - \mu)^0) \\ &= \sigma^2 \end{aligned}$$

μ & σ^2 unknown : $\hat{\sigma}^2 = s^2 = \frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n-1}$ is unbiased for σ^2

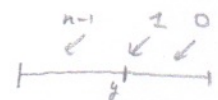
$$\begin{aligned} \text{(or) } s^2 &= \frac{1}{n-1} \sum_{i=1}^n [(X_i - \mu) + (\mu - \bar{X})]^2 \\ &= \frac{1}{n-1} \left[\sum_{i=1}^n (X_i - \mu)^2 - n(\bar{X} - \mu)^2 \right] \\ E_{\sigma^2}(s^2) &= \frac{1}{n-1} [n\sigma^2 - n\frac{\sigma^2}{n}] = \sigma^2 \checkmark \\ \text{(or) } E(\bar{X} - \mu)^2 &= \text{Var}(\bar{X}) = \frac{\sigma^2}{n} \end{aligned}$$

Ex 3 : $X_1, \dots, X_n$ IID Bernoulli(p) R.V. : let p be unknown

Prop : $\hat{p} = \frac{\sum_{i=1}^n X_i}{n}$ is unbiased for p (or) $E_p(\hat{p}) = \frac{1}{n} E(\sum X_i) = \frac{1}{n} (np) = p$

Ex 4 : $X_1, \dots, X_n$ IID $U(0, \theta)$ R.V. : let θ be unknown (so $\Theta = \{\theta > 0\}$)

Prop : $\hat{\theta} = 2\bar{X}$ is unbiased for θ (or) $E_{\theta}(\hat{\theta}) = 2E_{\theta}(\bar{X}) = 2\mu = \theta$ (or) $E_{\theta}(\bar{X}) = \mu$

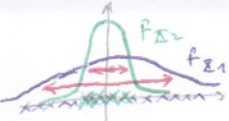


$\hat{\theta} = 2X_3$ is unbiased for θ (or) $E_{\theta}(\hat{\theta}) = 2E_{\theta}(X_3) = 2\mu = \theta$ (or) IID!

$\hat{\theta} = Y_n$ (n^{th} order stat) is biased for θ (or) $E(Y_n) = \int_0^{\theta} y \cdot f_{Y_n}(y) dy$ where $f_{Y_n}(y) = \frac{n!}{(n-1)!1!0!} \cdot (\frac{y}{\theta})^{n-1} \cdot (\frac{1}{\theta})^1 \cdot (1 - \frac{y}{\theta})^0$

$\hat{\theta} = \frac{n+1}{n} Y_n$ is unbiased for θ

$$\begin{aligned} &= \int_0^{\theta} \frac{n}{\theta} y y^{n-1} dy = \frac{n}{\theta} \int_0^{\theta} y^n dy \\ &= \frac{n}{\theta} \cdot \frac{y^{n+1}}{n+1} \Big|_0^{\theta} = \frac{n}{\theta} \cdot \frac{\theta^{n+1}}{n+1} = \frac{n}{n+1} \theta \rightarrow \text{correct it!} \\ f_{Y_n}(y) &= \begin{cases} \frac{n y^{n-1}}{\theta^n}, & 0 < y < \theta \\ 0, & \text{otherwise} \end{cases} \end{aligned}$$



Uniform Minimum-Variance Unbiased Estimators :

UMVUEs!

Uniform minimum variance unbiased estimators = (UMVUEs)

Let θ be an unknown parameter, $\theta \in \Theta$:

Then $\hat{\theta}$ is a UMVUE for θ if : • $E_{\theta}(\hat{\theta}) = \theta \quad \forall \theta \in \Theta$

• If $\hat{\theta}^*$ is another unbiased estimator for θ : $\text{Var}(\hat{\theta}) \leq \text{Var}(\hat{\theta}^*) \quad \forall \theta \in \Theta$

Note : UMVUEs do not always exist!

Prop = If a UMVUE exists, it is unique! $\rightarrow$ see HW2!

one way to find UMVUEs (when they exist) (another method later!)

Cramer-Rao Ineq =

$\rightarrow$ True also for PMF

$\tilde{X} = (X_1, \dots, X_n)$ a vector of r.v. with joint PDF : $f_{\tilde{X}}(\tilde{x}; \theta)$, depending on an unknown scalar $\theta \in \Theta$

then under certain regularity conditions :

if $E(T(\tilde{X})) = g(\theta)$, T a statistic, g a differentiable f^o

then
$$\text{Var}(T(\tilde{X})) \geq \frac{(g'(\theta))^2}{E\left[\left(\frac{\partial}{\partial \theta} \ln f_{\tilde{X}}(\tilde{X}, \theta)\right)^2\right]}$$

$\xrightarrow{\substack{\tilde{X}_i \text{ IID} \\ f_{\tilde{X}}(x, \theta) = \prod_{i=1}^n f_X(x_i, \theta)}} \frac{(g'(\theta))^2}{n E\left[\left(\frac{\partial}{\partial \theta} \ln f_X(X, \theta)\right)^2\right]}$

Note : • If $\hat{\theta}$ is an unbiased estimator for θ , then : if $\text{Var}(\hat{\theta}) = \text{C-R lower bound}$, $\hat{\theta}$ UMVUE!

• If $\hat{\theta}$ is unbiased, $g(\theta) = \theta \Rightarrow g'(\theta) = 1$

• C-R holds for PMFs $P_{\tilde{X}}$ also!

Corr : If $X_1, \dots, X_n$ are IID :
$$\text{Var}(T(\tilde{X})) \geq \frac{(g'(\theta))^2}{-n \cdot E\left[\frac{\partial^2}{\partial \theta^2} \ln f_{\tilde{X}}(\tilde{X}, \theta)\right]}$$
 (eq) HW2

Proof : (CR)

• claim : $E\left[\frac{\partial}{\partial \theta} \ln f_{\tilde{X}}(\tilde{X}, \theta)\right] \stackrel{\theta \in \Theta}{=} 0$ under the reg. cond^o

(*)
$$E\left[\frac{\partial}{\partial \theta} \ln f_{\tilde{X}}(\tilde{X}, \theta)\right] = \int_{-\infty}^{\infty} \frac{\partial}{\partial \theta} \ln f_{\tilde{X}}(\tilde{x}, \theta) \cdot f_{\tilde{X}}(\tilde{x}, \theta) d\tilde{x} = \int_{-\infty}^{\infty} \frac{f'_{\tilde{X}}(\tilde{x}, \theta)}{f_{\tilde{X}}(\tilde{x}, \theta)} f_{\tilde{X}}(\tilde{x}, \theta) d\tilde{x} = \frac{\partial}{\partial \theta} \int_{-\infty}^{\infty} f_{\tilde{X}}(\tilde{x}, \theta) d\tilde{x} = \frac{\partial}{\partial \theta} (1) = 0$$

• claim : $|g'(\theta)| \leq [E[(T(\tilde{X}) - g(\theta))^2]]^{1/2} [E\left[\left(\frac{\partial}{\partial \theta} \ln f_{\tilde{X}}(\tilde{X}, \theta)\right)^2\right]]^{1/2}$

(**)
$$|g'(\theta)| = \left| \frac{\partial}{\partial \theta} g(\theta) \right| = \left| \frac{\partial}{\partial \theta} E(T(\tilde{X})) \right| = \left| \frac{\partial}{\partial \theta} \int_{-\infty}^{\infty} T(\tilde{x}) \cdot f_{\tilde{X}}(\tilde{x}, \theta) d\tilde{x} \right| = \left| \int_{-\infty}^{\infty} T(\tilde{x}) \cdot \frac{\partial}{\partial \theta} f_{\tilde{X}}(\tilde{x}, \theta) d\tilde{x} \right|$$

$$\leq \int_{-\infty}^{\infty} |T(\tilde{x})| \cdot \left| \frac{\partial}{\partial \theta} \ln f_{\tilde{X}}(\tilde{x}, \theta) \right| \cdot f_{\tilde{X}}(\tilde{x}, \theta) d\tilde{x} \leq \int_{-\infty}^{\infty} |T(\tilde{x}) - g(\theta)| \cdot \left| \frac{\partial}{\partial \theta} \ln f_{\tilde{X}}(\tilde{x}, \theta) \right| \cdot f_{\tilde{X}}(\tilde{x}, \theta) d\tilde{x}$$

$\hookrightarrow$ add zero by prev claim!

(Cauchy-Schwartz)
$$\leq (E[(T(\tilde{X}) - g(\theta))^2])^{1/2} (E\left[\left(\frac{\partial}{\partial \theta} \ln f_{\tilde{X}}(\tilde{X}, \theta)\right)^2\right])^{1/2}$$

• $E(T(\tilde{X})) = g(\theta) \Rightarrow \text{Var}(T(\tilde{X})) = E[(T(\tilde{X}) - g(\theta))^2] \rightarrow$ Square both sides $\Rightarrow$ Get CR $\checkmark$

• By 1st claim : $E\left[\left(\frac{\partial}{\partial \theta} \ln f_{\tilde{X}}(\tilde{X}, \theta)\right)^2\right] \stackrel{(\dagger)}{=} \text{Var}\left[\frac{\partial}{\partial \theta} \ln f_{\tilde{X}}(\tilde{X}, \theta)\right] \stackrel{\tilde{X}_i \text{ indep}}{=} \text{Var}\left[\sum_{i=1}^n \frac{\partial}{\partial \theta} \ln f_{X_i}(\tilde{X}_i, \theta)\right]$

$$\stackrel{\substack{\uparrow i=1 \\ \tilde{X}_i \text{ indep}}}{=} \sum_{i=1}^n \text{Var}\left[\frac{\partial}{\partial \theta} \ln f_{X_i}(\tilde{X}_i, \theta)\right] = \sum_{i=1}^n E\left[\left(\frac{\partial}{\partial \theta} \ln f_{X_i}(\tilde{X}_i, \theta)\right)^2\right] = n E\left[\left(\frac{\partial}{\partial \theta} \ln f_X(X, \theta)\right)^2\right] \quad (*) \tilde{X}_i \text{ IID} \Rightarrow f_{\tilde{X}}(\tilde{x}) = \prod_{i=1}^n f_X(x_i) \text{ IID} \Rightarrow \ln f_{\tilde{X}}(\tilde{x}) = \sum_{i=1}^n \ln f_X(x_i) \Rightarrow \text{Var are equal!}$$

Examples:

- Bernoulli = $X_1, \dots, X_n$ IID Bernoulli (p) r.v.'s, p unknown: $\hat{p} = \frac{\sum_{i=1}^n X_i}{n} = \frac{\bar{X}}{n}$ UMVUE for p !

(8) Goal: $\text{Var}[\hat{p}] = \text{CR lower bound!}$

$$\text{Var}[\hat{p}] = \text{Var}\left(\frac{\bar{X}}{n}\right) = \frac{n p(1-p)}{n^2} = \frac{p(1-p)}{n}$$

(8) Indep $\Rightarrow \text{Var } \bar{X} = \sum \text{Var}$
Id. dist. $\Rightarrow p(1-p)$

• CR Lower bound:

$$\begin{aligned} n E\left[\left(\frac{\partial}{\partial p} \ln p_X(\bar{X}, p)\right)^2\right] &= n E\left[\left(\frac{\partial}{\partial p} \ln(p^{\bar{X}}(1-p)^{n-\bar{X}})\right)^2\right] = n E\left[\left(\frac{\partial}{\partial p} (\bar{X} \ln p + (n-\bar{X}) \ln(1-p))\right)^2\right] \\ &= n E\left[\left(\frac{\bar{X}}{p} - \frac{n-\bar{X}}{1-p}\right)^2\right] = n E\left[\left(\frac{\bar{X}-p}{p(1-p)}\right)^2\right] = \frac{n}{(p(1-p))^2} E[(\bar{X}-p)^2] = \frac{n}{(p(1-p))^2} \text{Var}(\bar{X}) = \frac{n}{(p(1-p))^2} \frac{p(1-p)}{n} \\ &= \frac{1}{p(1-p)} = \frac{1}{\text{Var}[\hat{p}]} \end{aligned}$$

$\Rightarrow$ CR lower bound = $\text{Var}[\hat{p}]$ ✓

- Normal = $X_1, \dots, X_n$ IID $N(\mu, \sigma^2)$ r.v.'s:

• μ unknown: $\hat{\mu} = \bar{X}$ is UMVUE for μ (True for σ^2 known, and for σ^2 unknown)

• σ^2 unknown: $\rightarrow \mu = \mu_0$ known: $\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \mu_0)^2$
 $\rightarrow \mu$ unknown: $\hat{S}^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$ } UMVUEs for σ^2

- Uniform = $X_1, \dots, X_n$ IID $U(0, \theta)$ r.v.'s, θ unknown:

• $\hat{\theta} = 2\bar{X}$ and $\hat{\theta} = \frac{n+1}{n} Y_n$ are unbiased estimators for θ

• $\hat{\theta} = \frac{n+1}{n} Y_n$ is the UMVUE $\rightarrow \text{Var}(2\bar{X}) > \text{Var}(\frac{n+1}{n} Y_n)$ $\triangle$ Cannot use C-R: regularity cond^o fail!
 (8) range of PDF depends on the param θ .

Note: • Some UMVUEs do not reach the CR lower bound, even when reg cond^o satisfied

• UMVUEs depend on the distribution family (Ex: S^2 UMVUE for $X_i \sim N(\mu, \sigma^2)$
 S^2 NOT UMVUE for $\beta^2 = \text{Var}(X)$, $X_i \sim \text{Exp}(\beta)$)

Mean-Square Error = MSE is the "Average" squared distance between $\hat{\theta}$ and θ :

$$\hat{\theta} \text{ an estimator for } \theta: \text{MSE}(\hat{\theta}) = E[(\hat{\theta} - \theta)^2]$$

$$\text{MSE}(\hat{\theta}) = \text{Var}(\hat{\theta}) + (\text{Bias})^2 \quad \text{where Bias} = (E(\hat{\theta}) - \theta)$$

$$\begin{aligned} \text{Proof: } \text{MSE}(\hat{\theta}) &= E[(\hat{\theta} - \theta)^2] = E[(\hat{\theta} - E[\hat{\theta}] + E[\hat{\theta}] - \theta)^2] \\ &= \text{Var}(\hat{\theta}) + E[(\hat{\theta} - E[\hat{\theta}]) \cdot (E[\hat{\theta}] - \theta)] + E[(E[\hat{\theta}] - \theta)^2] \\ &= 0 \quad (\text{8}) = \underbrace{E(\hat{\theta} - \theta)}_{\text{EST}} \cdot \underbrace{E(\hat{\theta} - E(\hat{\theta}))}_{=0} + E[(E[\hat{\theta}] - \theta)^2] \end{aligned}$$

$$\text{Note} = \left\{ \begin{aligned} S^2 &= \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 \rightarrow \text{Bias} = 0; \text{ larger MSE!} \\ \sigma^2 &= \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \rightarrow \text{Bias} > 0; \text{ smaller MSE!} \end{aligned} \right\} \rightarrow \text{under Normality!}$$

Consistency :

As Sample size $n \rightarrow \infty : \hat{\theta} \xrightarrow{P} \theta$

Consistency = An estimator $\hat{\theta}$ for θ is consistent for θ if : $\hat{\theta} \xrightarrow{P} \theta$ as $n \rightarrow \infty$

i.e. : $P(|\hat{\theta} - \theta| > \epsilon) \rightarrow 0$ as $n \rightarrow \infty$

i.e. : $P(|\hat{\theta} - \theta| < \epsilon) \rightarrow 1$ as $n \rightarrow \infty$ $\forall \epsilon > 0$.

Methods for establishing consistency =

1] Try to use the Weak law of large nb (WLLN):

WLLN = $X_1, \dots, X_n$ IID r.v. with finite mean μ : Then $\frac{1}{n} \sum_{i=1}^n X_i \xrightarrow{P} \mu$

Method = • Try to write $\hat{\theta}(\tilde{X})$ as $\hat{\theta}(\bar{X})$ for some $f = \hat{\theta}$ of $\bar{X}$.

• Then $g(\bar{X}) \xrightarrow{P} g(\mu)$ $\because \bar{X} \xrightarrow{P} \mu$

• If $g(\mu) = \theta$, then we are done!

2] If : $\begin{cases} E(\hat{\theta}_n) \rightarrow \theta \\ \text{Var}(\hat{\theta}_n) \rightarrow 0 \end{cases}$ as $n \rightarrow \infty$, then $\hat{\theta}_n \xrightarrow{P} \theta$

3] Directly : use the def of convergence in probability

Examples =

• $\bar{X}$ is consistent for μ $(\because)$ WLLN

• S^2 is consistent for σ^2 (if $E(X^2) < \infty$) $(\because)$

$$S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{1}{n-1} \left[\sum_{i=1}^n X_i^2 - n\bar{X}^2 \right] = \frac{n}{n-1} \left[\frac{1}{n} \sum_{i=1}^n X_i^2 - \bar{X}^2 \right]$$

$$\xrightarrow{S_i \text{ IID}} \xrightarrow{S_i^2 \text{ IID}} E(S_i^2) < \infty : WLLN \Rightarrow \frac{1}{n} \sum_{i=1}^n X_i^2 \xrightarrow{P} E(X^2) = \sigma^2 + \mu^2$$

$$\Rightarrow S^2 \xrightarrow{P} \frac{n}{n-1} [\sigma^2 + \mu^2 - \mu^2] \rightarrow \sigma^2 \text{ as } n \rightarrow \infty$$

• $X_1, \dots, X_n$ IID Bernoulli(p) : $\hat{p}$ is consistent for p

$$\hat{p} = \frac{1}{n} \sum_{i=1}^n X_i$$

Sample proportion (prob of Success)

$(\because)$ WLLN : replace μ by " p ".

• $X_1, \dots, X_n$ IID Exp(β) : $e^{-1/\bar{X}}$ is consistent for $P(X_1 > 1) = e^{-1/\beta}$

Note : $\text{Var}(X_i) = \beta^2 \Rightarrow \begin{cases} \bar{X}^2 \xrightarrow{P} \beta^2 \\ S^2 \xrightarrow{P} \beta^2 \end{cases} \rightarrow \text{has min Variance!}$

$(\because)$ WLLN : $\bar{X} \xrightarrow{P} E(X_1) = \beta$
 $x \mapsto e^{-1/x} \text{ cont} \Rightarrow e^{-1/\bar{X}} \xrightarrow{P} e^{-1/\beta}$

• $X_1, \dots, X_n$ IID $U(0, \theta)$: $\begin{cases} 2\bar{X} \xrightarrow{P} 2\mu = \theta \\ \frac{n+1}{n} Y_n \xrightarrow{P} \theta \end{cases}$ are both consistent for θ !

$(\because) 2\bar{X} \rightarrow 2\mu = 2 \cdot \frac{\theta}{2} = \theta$

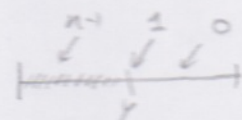
Proof : • direct : Note : $\frac{n+1}{n} Y_n$ consistent for $\theta \Leftrightarrow Y_n$ consistent for θ . $(\because) \frac{n+1}{n} \rightarrow 1$

Trick : Given $\epsilon > 0$: $P[|Y_n - \theta| > \epsilon] = P[|\theta - Y_n| > \epsilon] = P[\theta - Y_n > \epsilon] = P[Y_n < \theta - \epsilon]$
 $P(Y_n < \theta - \epsilon) = P(X_1 < \theta - \epsilon, \dots, X_n < \theta - \epsilon)$ $P(Y_n > \theta) = 0$
 $[X_i \text{ IID}] = [P(X_1 < \theta - \epsilon)]^n = (F_{X_1}(\theta - \epsilon))^n = \left(\frac{\theta - \epsilon}{\theta}\right)^n = \left(1 - \frac{\epsilon}{\theta}\right)^n = 0 \text{ if } \epsilon > \theta$

• Alternate :

• $E(Y_n) \rightarrow \theta$ $(\because) E(Y_n) = \left(\frac{n}{n+1}\right) \theta \rightarrow \theta$

• $\text{Var}(Y_n) \rightarrow 0$ $(\because) f_{Y_n}(y) = \frac{n!}{1!0!(n-1)!} \left(\frac{y}{\theta}\right)^{n-1} \cdot \frac{1}{\theta} = \frac{n y^{n-1}}{\theta^n}$ for $0 < y < \theta$



$$\Rightarrow \text{Var}(Y_n) = E(Y_n^2) - [E(Y_n)]^2 = \int_0^\theta y^2 \frac{n y^{n-1}}{\theta^n} dy - \left(\frac{n}{n+1}\theta\right)^2 = \left(\frac{n}{n+2}\right)\theta^2 - \left(\frac{n}{n+1}\right)^2\theta^2 \rightarrow 0$$

Sufficiency

→ Once you have the value of a Sufficient Statistic, you have extracted all the information contained in the observations about the parameter θ .

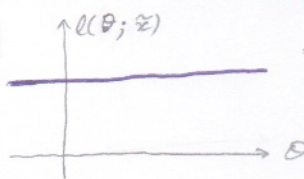
① Definition

Likelihood function: Let $f_{\tilde{X}}(\tilde{x}; \theta)$ be the PDF of $\tilde{X}$ at θ fixed
 $l(\theta; \tilde{x})$ is the function of θ at $\tilde{x}$ fixed $\Rightarrow$ "likelihood function"

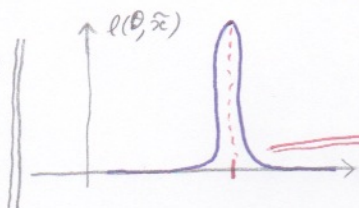
Note: • To compute $l(\theta; \tilde{x})$, you need $f_{\tilde{X}}(\tilde{x}; \theta)$ for θ and $\tilde{x}$.
 • All the info contained in the data $\tilde{x}$ (fixed) about θ is captured by $l(\theta; \tilde{x})$.

→ How the likelihood changes with θ !

Ex:



→ Not very informative about θ :
 $l(\theta; \tilde{x})$ varies little at θ
 for any given $\tilde{x}$.



$\Rightarrow$ Informative about θ :
 $l(\theta; \tilde{x})$ changes markedly with θ .
 gives the values of θ that are more plausible to have led to the data $\tilde{x}$.

Sufficiency = Let $\tilde{X}$ have joint PDF $f_{\tilde{X}}(\tilde{x}; \theta)$

A statistic T is Sufficient for θ if $f_{\tilde{X}|T}(\tilde{x}; \theta | T(\tilde{x}) = t)$ is indep of $\theta \forall \tilde{x}, \forall t$

→ once you know $t = t(\tilde{x})$ for every $\tilde{x}$, there is no info left in the data about θ .

Ex = $X_1, \dots, X_n$ IID Bernoulli(p) r.v., p unknown : $T(\tilde{X}) = \sum_{i=1}^n X_i$ is sufficient for p ($T = \#$ of Successes)

Goal: Fix $\tilde{x}$ and $\tilde{x}$: $P_{\tilde{X}|T}(\tilde{x}; p | T(\tilde{x}) = t)$ indep of $p \forall t$

$$(\text{or}) P_{\tilde{X}|T}(\tilde{x}; p | T(\tilde{x}) = t) = P(X_1 = x_1, \dots, X_n = x_n; p | \sum_{i=1}^n X_i = t)$$

• if $t \neq \sum_{i=1}^n x_i$, then $\sum_{i=1}^n X_i \neq t \Rightarrow P_{\tilde{X}|T}(\tilde{x}; p | T(\tilde{x}) = t) = 0 \Rightarrow$ indep of p ✓

• if $t = \sum_{i=1}^n x_i$: $P_{\tilde{X}|T}(\tilde{x}; p) = \frac{P(X_1 = x_1, \dots, X_n = x_n; p | T(\tilde{x}) = t)}{P(T(\tilde{x}) = t)} = \frac{P(X_1 = x_1, \dots, X_n = x_n; p)}{P(\sum_{i=1}^n X_i = t)} = \frac{\text{IID Ber}(p)}{\text{Binomial: } t \text{ Success in } n \text{ indep trials}}$

$$\Rightarrow P_{\tilde{X}|T}(\tilde{x}; p) = \frac{p^t (1-p)^{n-t}}{\binom{n}{t} p^t (1-p)^{n-t}} = \frac{1}{\binom{n}{t}} \text{ for } t=0, 1, \dots, n \Rightarrow \text{indep of } p \text{ for every } \tilde{x} \quad \checkmark$$

Note = • A Sufficient statistic can be a vector $\tilde{T} = (T_1, \dots, T_n)$ $\text{Ex: } \tilde{T} = (\sum_{i=1}^n X_i, \sum_{i=1}^n (X_i - \bar{X})^2)$

• Trivial Sufficient statistic: $\tilde{T}(\tilde{X}) = \tilde{X} \rightarrow$ always sufficient for θ .
 $= (X_1, X_2, \dots, X_n)$

• Let T a sufficient statistic, g a 1-1 function: $T' = g(T)$ is also Sufficient
 (or) given $T=t$ indep of $\theta \Rightarrow g(T)=g(t)$ indep of θ as g is 1-1.

• Sufficient statistics may not be "good" estimators of θ (Ex: $\tilde{X} = \text{Ber}(p)$, p unknown:
 $T = \sum_{i=1}^n X_i$ & $T' = \frac{1}{n} \sum_{i=1}^n X_i$ Sufficient for p , but T' is UMVUE, T is bad + consistent)

• Sufficiency is a family property:
 T sufficient for a family of dist $\not\Rightarrow T$ sufficient for another family

Ex: $X_1, \dots, X_n$ IID $U(0, \theta)$, θ unknown

$$T(\tilde{X}) = \bar{X} \text{ is NOT sufficient for } \theta: f_{\tilde{X}|\Sigma \tilde{X}}(\tilde{x}; \theta | \sum \tilde{X}_i = t) = \frac{f_{\tilde{X}}(\tilde{x}; \theta)}{\sum_{\tilde{X}(\tilde{x})=t} f_{\tilde{X}}(\tilde{x}; \theta)} = \dots \text{ depends on } \theta!$$

• If the data $\tilde{X} = (x_1, \dots, x_n)$ is destroyed and $T(\tilde{X})$ sufficient is known: we can generate data consistent with $\tilde{x}$ but we cannot recover $\tilde{x}$

Neymann-Fisher Factorization Theorem:

N-F Thm: A statistic T is sufficient $\iff$ we can write $f_{\tilde{x}}(\tilde{x}; \theta) = g(T(\tilde{x}); \theta) h(\tilde{x})$

where:
 • g depends on θ and on the observations $\tilde{x}$ only through t .
 • h is a function of the observations $\tilde{x}$ only (h indep of θ !)

Proof = (the proof is done in the discrete case: continuous case true but proof more technical!)

For $\tilde{x}$ discrete:

$$\begin{aligned} \Rightarrow \text{Let } \tilde{x} = \tilde{x}_0: P_{\tilde{x}}(\tilde{x}_0; \theta) &= P(\tilde{X} = \tilde{x}_0; \theta) = P(\tilde{X} = \tilde{x}_0, T(\tilde{X}) = T(\tilde{x}_0); \theta) = \underbrace{P(\tilde{X} = \tilde{x}_0; \theta | T(\tilde{X}) = T(\tilde{x}_0))}_{h(\tilde{x}_0)} \cdot \underbrace{P(T(\tilde{X}) = T(\tilde{x}_0); \theta)}_{g(T(\tilde{x}_0); \theta)} \\ \Rightarrow \text{By assumption: } P(\tilde{X} = \tilde{x}_0; \theta) &= g(T(\tilde{x}_0); \theta) h(\tilde{x}_0) \\ \text{So: } P(\tilde{X} = \tilde{x}_0; \theta | T(\tilde{X}) = T(\tilde{x}_0)) &= \frac{P(\tilde{X} = \tilde{x}_0; \theta)}{P(T(\tilde{X}) = T(\tilde{x}_0))} = \frac{g(T(\tilde{x}_0); \theta) h(\tilde{x}_0)}{\sum_{\tilde{x}: T(\tilde{x}) = T(\tilde{x}_0)} P(\tilde{X} = \tilde{x}; \theta)} = \frac{g(T(\tilde{x}_0); \theta) h(\tilde{x}_0)}{\sum_{\tilde{x}: T(\tilde{x}) = T(\tilde{x}_0)} g(T(\tilde{x}); \theta) h(\tilde{x})} = \frac{g(T(\tilde{x}_0); \theta) h(\tilde{x}_0)}{\sum_{\tilde{x}: T(\tilde{x}) = T(\tilde{x}_0)} g(T(\tilde{x}); \theta) h(\tilde{x})} \\ \Rightarrow P(\tilde{X} = \tilde{x}_0; \theta | T(\tilde{X}) = T(\tilde{x}_0)) &= \frac{h(\tilde{x}_0)}{\sum_{\tilde{x}: T(\tilde{x}) = T(\tilde{x}_0)} h(\tilde{x})} \quad \text{indep of } \theta \forall \tilde{x} \Rightarrow T \text{ sufficient for } \theta! \end{aligned}$$

Examples =

① $X_1, \dots, X_n$ IID Bernoulli(p) r.v., p unknown: $\leadsto T(\tilde{X}) = \sum_{i=1}^n X_i$ sufficient for p .

$$\begin{aligned} \text{(*) } P(X_1 = x_1, \dots, X_n = x_n; p) &= p^{\sum_{i=1}^n x_i} (1-p)^{n - \sum_{i=1}^n x_i} = \underbrace{p^{\sum_{i=1}^n x_i}}_{g(\sum_{i=1}^n x_i; p)} \cdot \underbrace{\prod_{i=1}^n I[x_i \in \{0, 1\}]}_{h(\tilde{x})} \xrightarrow{\text{N-F Thm}} T(\tilde{X}) = \sum_{i=1}^n X_i \text{ sufficient for } p. \\ &\quad \uparrow \text{True } \forall p \end{aligned}$$

② $X_1, \dots, X_n$ IID $U(0, \theta)$ r.v., θ unknown: $\leadsto T(\tilde{X}) = \max\{X_1, \dots, X_n\}$ is sufficient for θ .

$$\text{(*) } f_{\tilde{x}}(\tilde{x}; \theta) = \frac{1}{\theta^n} \prod_{i=1}^n I[0 < x_i < \theta] = \frac{1}{\theta^n} \prod_{i=1}^n I[x_i > 0] \cdot \prod_{i=1}^n I[x_i < \theta] = \underbrace{\frac{1}{\theta^n} \cdot I[\max_{1 \leq i \leq n} x_i < \theta]}_{g(T(\tilde{x}); \theta)} \cdot \underbrace{I[\min_{1 \leq i \leq n} x_i > 0]}_{h(\tilde{x})} \quad \forall \tilde{x}, \forall \theta$$

③ $X_1, \dots, X_n$ IID Gamma(α, β) r.v.:
 $\alpha = ?$, β known: $\leadsto T(\tilde{X}) = \sum_{i=1}^n X_i$ sufficient for α
 α known, $\beta = ?$: $\leadsto T(\tilde{X}) = \sum_{i=1}^n X_i$ sufficient for β
 α, β unknown: $\leadsto \tilde{T}(\tilde{X}) = (\sum_{i=1}^n X_i, \sum_{i=1}^n 1/X_i)$ sufficient for (α, β) .

$$\text{(*) } f_{\tilde{x}}(\tilde{x}; \alpha, \beta) = \frac{1}{(\Gamma(\alpha))^n} \cdot \frac{1}{\beta^n} \cdot \prod_{i=1}^n x_i^{\alpha-1} \cdot e^{-\sum_{i=1}^n x_i/\beta} \cdot \prod_{i=1}^n I[0 < x_i < \infty] \Rightarrow \text{arrange terms for each case.}$$

$$\text{Ex: Case 1: } f_{\tilde{x}}(\tilde{x}; \alpha, \beta) = \underbrace{\left(\prod_{i=1}^n x_i \right)^{\alpha-1} \cdot \frac{1}{(\Gamma(\alpha))^n}}_{g(\sum_{i=1}^n x_i, \alpha)} \cdot \underbrace{\frac{1}{\beta^n} e^{-\sum_{i=1}^n x_i/\beta}}_{h(\tilde{x})} \cdot \underbrace{\prod_{i=1}^n I[0 < x_i < \infty]}_{\text{for } \beta \text{ known}}$$

Note: even if T is sufficient, you still want to keep your data to check your model (IID Gamma, Binom, Bernoulli, Normal, ...)

Indicator function: $I[A] = \begin{cases} 1, & x \in A \\ 0, & x \notin A \end{cases}$

Minimal Sufficiency: → Best Summarization of the data $\mathbf{x}_1, \dots, \mathbf{x}_n$

minimal Sufficient Statistics = a statistic T is minimal sufficient if

- (1) T is sufficient for Θ
- (2) T is a function of any other sufficient statistic T' : if T' is sufficient, then $\exists g$ s.t. $T = g(T')$

Lemma = $\tilde{x}, \tilde{y}$ two observations from $f_{\tilde{x}}(\cdot; \theta)$, T a sufficient statistic:

$$T(\tilde{x}) = T(\tilde{y}) \Rightarrow \frac{f_{\tilde{x}}(\tilde{x}; \theta)}{f_{\tilde{y}}(\tilde{y}; \theta)} \text{ is independent of } \theta.$$

Proof: NFHM $\Rightarrow \frac{f_{\tilde{x}}(\tilde{x}; \theta)}{f_{\tilde{y}}(\tilde{y}; \theta)} = \frac{g(T(\tilde{x}), \theta) h(\tilde{x})}{g(T(\tilde{y}), \theta) h(\tilde{y})} \stackrel{T(\tilde{x})=T(\tilde{y})}{=} \frac{h(\tilde{x})}{h(\tilde{y})}$

Theorem = T a statistic: if " $\frac{f_{\tilde{x}}(\tilde{x}; \theta)}{f_{\tilde{y}}(\tilde{y}; \theta)}$ is independent of $\theta \Leftrightarrow T(\tilde{x}) = T(\tilde{y})$ " holds, then T is minimal sufficient for Θ .

Proof = Let $\tilde{X}$ be discrete ($\tilde{X}$ continuous also works but technical!)

• (sufficiency): $P(\tilde{X} = \tilde{x}_0; \theta | T(\tilde{X}) = T(\tilde{x}_0)) = \frac{P(\tilde{X} = \tilde{x}_0; \theta)}{\sum_{\tilde{x}: T(\tilde{x})=T(\tilde{x}_0)} P(\tilde{X} = \tilde{x}; \theta)} = \frac{1}{\sum_{\tilde{x}: T(\tilde{x})=T(\tilde{x}_0)} \frac{P(\tilde{X} = \tilde{x}; \theta)}{P(\tilde{X} = \tilde{x}_0; \theta)}} = \text{indep of } \theta.$

• (minimality): We know that T is sufficient: let S be any other sufficient statistic:

S is sufficient: $\forall \tilde{w}, \tilde{z}: S(\tilde{w}) = S(\tilde{z}) \xRightarrow{\text{Lemma}} \frac{P_{\tilde{X}}(\tilde{w}; \theta)}{P_{\tilde{X}}(\tilde{z}; \theta)} \text{ indep of } \theta \xRightarrow{\text{"hypothesis"}} T(\tilde{w}) = T(\tilde{z})$

so $\forall \tilde{w}, \tilde{z}: S(\tilde{w}) = S(\tilde{z}) \Rightarrow T(\tilde{w}) = T(\tilde{z})$, hence T is a function of S , so T minimal sufficient.

Examples =

① $X_1, \dots, X_n$ IID Poisson(λ) r.v., λ unknown: $T(\tilde{X}) = \sum_{i=1}^n X_i$ is minimal sufficient for λ

(*) $\frac{P(\tilde{X} = \tilde{x}; \lambda)}{P(\tilde{X} = \tilde{y}; \lambda)} = \frac{\lambda^{\sum x_i} e^{-\lambda} / \prod x_i!}{\lambda^{\sum y_i} e^{-\lambda} / \prod y_i!} = \left(\prod \frac{x_i!}{y_i!} \right) \lambda^{\sum (x_i - y_i)}$ is indep of $\lambda \Leftrightarrow \sum x_i = \sum y_i$ ✓ so $T = \sum X_i$ is min. Suff.

② $X_1, \dots, X_n$ IID $N(\mu, \sigma^2)$, $\mu \neq 0$ known, σ^2 unknown: $T(\tilde{X}) = \sum_{i=1}^n (X_i - \mu)^2$ is min. Suff. for σ^2

(*) $\frac{f_{\tilde{X}}(\tilde{x}; \mu, \sigma^2)}{f_{\tilde{Y}}(\tilde{y}; \mu, \sigma^2)} = \frac{\left(\frac{1}{\sigma\sqrt{2\pi}}\right)^n e^{-\frac{1}{2\sigma^2} \sum (x_i - \mu)^2}}{\left(\frac{1}{\sigma\sqrt{2\pi}}\right)^n e^{-\frac{1}{2\sigma^2} \sum (y_i - \mu)^2}} = e^{-\frac{1}{2\sigma^2} [\sum (x_i - \mu)^2 - \sum (y_i - \mu)^2]}$ is indep of $\sigma^2 \Leftrightarrow \sum (x_i - \mu)^2 = \sum (y_i - \mu)^2$

Note: if $\mu = 0$, $T(\tilde{X}) = \sum_{i=1}^n X_i^2$ is min. Suff for σ^2 .

Completeness :

→ How to find UMVUES

① Definitions :

Completeness of a family = A family of distribution $\{f_{\tilde{X}}(\tilde{x}, \theta) : \theta \in \Theta\}$ is complete

$$\Leftrightarrow "E_{\theta}[g(\tilde{X})] \stackrel{?}{=} 0 \Rightarrow g(\tilde{X}) = 0 \text{ with prob } 1" \text{ holds } \forall \text{ measurable fnd } g.$$

Completeness of a statistic = A statistic T is complete if the family of distributions induced by T is complete.

$$\text{i.e. } "E_{\theta}[g(T(\tilde{X}))] \stackrel{?}{=} 0 \Rightarrow g(T(\tilde{X})) = 0 \text{ with prob } 1" \text{ holds } \forall \text{ meas. fnd } g.$$

Examples =

① $X \sim \text{Poisson}(\lambda)$ is a complete family:

$$(\otimes) E_{\lambda}[g(X)] \stackrel{?}{=} 0 \Leftrightarrow \sum_{x=0}^{\infty} g(x) \frac{\lambda^x e^{-\lambda}}{x!} \stackrel{?}{=} 0 \Leftrightarrow \sum_{x=0}^{\infty} \frac{g(x)}{x!} \lambda^x \stackrel{?}{=} 0 \xrightarrow{\text{power series}} g(x) = 0 \forall x = 0, 1, \dots, \infty$$

but in Poisson: $P(X \in \mathbb{N}) = 1$, so $g(x) = 0 \forall x \in \mathbb{R}$ with probability 1 ✓

② $X_1, \dots, X_n \text{ IID Poisson}(\lambda)$ r.v. : $T = \sum_{i=1}^n X_i$ is complete.

$$(\otimes) \sum_{i=1}^n X_i \sim \text{Poisson}(n\lambda), \text{ and Poisson family is complete} \Rightarrow T \text{ is complete!}$$

③ $X_1, \dots, X_n \text{ IID } U(0, \theta)$: $T = Y_n$ (max order statistic) is complete.

$$(\otimes) f_{Y_n}(y) = \frac{n y^{n-1}}{\theta^n} \mathbb{I}[0 < y < \theta], \text{ so } E_{\theta}[g(Y_n)] = \int_0^{\theta} g(y) \frac{n y^{n-1}}{\theta^n} dy \stackrel{?}{=} 0 \Leftrightarrow \int_0^{\theta} g(y) y^{n-1} dy = 0 \forall \theta > 0$$

$$\Rightarrow \frac{d}{dy}(\dots) \stackrel{?}{=} 0 \Rightarrow g(\theta) \theta^{n-1} \stackrel{?}{=} 0 \Rightarrow g(x) = 0 \forall 0 < x < \theta$$

$$\Rightarrow g(\tilde{X}) = 0 \text{ with prob } 1$$

Rao-Blackwell improvement thm =

Let $\hat{\theta}$ an unbiased estimator for θ ,
 T a Sufficient statistic for θ , define $T' := E(\hat{\theta} | T)$

Then: T' is unbiased for θ ($E_{\theta}(T') = \theta$) and $\text{Var}_{\theta}(T') \leq \text{Var}_{\theta}(\hat{\theta})$

Proof = Recall: $E[\hat{\theta} | T] := E[\hat{\theta} | T=t]$ is a r.v., as it varies randomly with T .

- T' is an estimator for θ ($\otimes$) T is sufficient for θ , so T' does not depend on θ .
- $E[T'] = E_T[E_{\theta}[\hat{\theta} | T]] = E_{\theta}(\hat{\theta}) = \theta$ ($\otimes$) $\hat{\theta}$ unbiased by assumption.
- $\text{Var}(\hat{\theta}) = \underbrace{E[\text{Var}(\hat{\theta} | T)]}_{\geq 0} + \text{Var}[E(\hat{\theta} | T)] \geq \text{Var}[E(\hat{\theta} | T)] = \text{Var}(T')$

Note = Even if $\hat{\theta}$ is not unbiased for θ , R-B thm will improve the Mean Square Error of $\hat{\theta}$.

$$\text{i.e. } \text{MSE}(\hat{\theta}) \geq \text{MSE}(T')$$

($\otimes$) $\hat{\theta}$ given, T sufficient, $T' := E(\hat{\theta} | T)$:

$$\begin{aligned} \text{MSE}(\hat{\theta}) &= E(\hat{\theta} - \theta)^2 = \text{Var}(\hat{\theta}) + (E[\hat{\theta} - \theta])^2 \geq \text{Var}(T') + (E(\hat{\theta}) - \theta)^2 = \text{Var}(T') + (E_T[E(\hat{\theta} | T)] - \theta)^2 \\ &= \text{Var}(T') + (E[T'] - \theta)^2 \\ &= \text{Var}(T') + (E[T' - \theta])^2 \\ &= \text{MSE}(T') \end{aligned}$$

IV Lehmann - Scheffé Theorem: → Link between Sufficiency, Completeness & UMVUEs

L-S thm = $\hat{\theta}$ an unbiased estimator for θ ; T a complete & sufficient statistic:

If $\exists g$ s.t. $\hat{\theta} = g(T)$, then $\hat{\theta}$ is the unique UMVUE for θ .

Proof = • Uniqueness: Suppose $\exists$ another unbiased $g'(T)$: claim: $g(T) = g'(T)$ with prob 1.

$$\begin{aligned} \text{(\&)} \quad \left. \begin{aligned} E(g(T)) &= E(\hat{\theta}) \stackrel{!}{=} \theta \\ E(g'(T)) &= E(\hat{\theta}) \stackrel{!}{=} \theta \end{aligned} \right\} &\Rightarrow E(g(T)) - E(g'(T)) \stackrel{!}{=} 0 &\Rightarrow E[g(T) - g'(T)] \stackrel{!}{=} 0 \\ &&&\Rightarrow E[(g - g')(T)] \stackrel{!}{=} 0 \Rightarrow (g - g')(T) = 0 \text{ with prob 1} \end{aligned}$$

• UMVUE: Suppose $\hat{\theta}'$ is another unbiased estimator for θ with $\text{Var}(\hat{\theta}') < \text{Var}(\hat{\theta})$

Then improve $\hat{\theta}'$ by Rao-Blackwellizing it ($\because T$ sufficient): $g'(T) = E(\hat{\theta}'|T)$ unbiased.

$$\text{So } \text{Var}(g'(T)) \leq \text{Var}(\hat{\theta}') < \text{Var}(\hat{\theta}) = \text{Var}(g(T))$$

$$\text{So } \text{Var}(g'(T)) < \text{Var}(g(T)) \quad \Rightarrow \quad \text{(\&)} \quad g(T), g'(T) \text{ are unbiased estimators of } \theta, \& \text{ functions of } T$$

So by "claim" above: $g(T) = g'(T)$ with prob 1.

How to apply the L-S thm:

① Hope to find an unbiased estimator $\hat{\theta}$ of θ , s.t. $\hat{\theta} = g(T)$ where T is a complete sufficient statistic.

OR

② If you have: • an unbiased estimator $\hat{\theta}$ of θ
• a complete sufficient statistic T } Rao-Blackwellize $\hat{\theta}$ by conditioning on T :
 $T' := E[\hat{\theta}|T]$ is UMVUE!

Tip: take T as simple as you can!

Examples =

① $X_1, \dots, X_n$ IID Poisson(λ) r.v., λ unknown: $\begin{cases} \hat{\lambda} = \bar{X} \text{ is UMVUE for } \lambda \\ \hat{\theta} = (1 - \frac{1}{n})^{\sum_{i=1}^n X_i} \text{ is UMVUE for } p(X=0) = e^{-\lambda} \end{cases}$

(\&) • $\hat{\lambda} = \bar{X}$ is unbiased for λ , and $T = \sum_{i=1}^n X_i$ is complete + sufficient for $\lambda \xrightarrow{\text{L-S thm}} \hat{\lambda} = g(T) = \frac{1}{n}T$ is UMVUE for λ .

• Simple unbiased estimator $\hat{\theta}$ of $e^{-\lambda}$: $\hat{\theta} := I[X=0]$ (\&) $E[I[X=0]] = P(X=0) = e^{-\lambda}$.

$T = \sum_{i=1}^n X_i$ is complete + sufficient $\Rightarrow T' = E[I[X_1=0] | \sum_{i=1}^n X_i = t]$ is UMVUE for $e^{-\lambda}$.

$$\Rightarrow T' = \frac{e^{-\lambda} \frac{[(n-1)\lambda]^t e^{-(n-1)\lambda}}{(n\lambda)^t e^{-n\lambda} / t!}}{(n\lambda)^t e^{-n\lambda} / t!} = \left(\frac{n-1}{n} \right)^t \Big|_{t=T} \Rightarrow T' = \left(1 - \frac{1}{n} \right)^T = \left(1 - \frac{1}{n} \right)^{\sum_{i=1}^n X_i}$$

$$\text{Note} = \left(1 - \frac{1}{n} \right)^{\sum_{i=1}^n X_i} = \left(1 - \frac{1}{n} \right)^{n\bar{X}} \xrightarrow{n \rightarrow \infty} (e^{-1})^{\lambda} \text{ as } \bar{X} \xrightarrow{n \rightarrow \infty} \lambda \text{ and } \left(1 - \frac{1}{n} \right)^n \xrightarrow{n \rightarrow \infty} e^{-1}$$

② $X_1, \dots, X_n$ IID $U(0, \theta)$ r.v., θ unknown: $\hat{\theta}_n = Y_n$ is UMVUE for θ
(\&) $\hat{\theta}_n = \frac{n+1}{n} Y_n$ unbiased for θ , $T = Y_n$ is complete & sufficient. $\xrightarrow{\text{L-S thm}} \hat{\theta} = g(Y_n) = \frac{n+1}{n} Y_n$ is UMVUE for θ .

③ $N(\mu, \sigma^2)$, μ unknown, σ^2 known forms a complete family
(\&) $E_{\theta}[g(x)] \stackrel{!}{=} 0 \Leftrightarrow \int_{-\infty}^{\infty} g(x) \cdot \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}} dx \stackrel{!}{=} 0 \Leftrightarrow \int_{-\infty}^{\infty} g(x) e^{-\frac{x^2 - 2\mu x}{2\sigma^2}} dx \stackrel{!}{=} 0 \Leftrightarrow \int_{-\infty}^{\infty} \underbrace{g(x) e^{-\frac{x^2}{2\sigma^2}}}_{=: h(x)} e^{-\frac{\mu}{\sigma^2} x} dx \stackrel{!}{=} 0$
 $\Rightarrow \int_{-\infty}^{\infty} h(x) e^{-\frac{\mu}{\sigma^2} x} dx \stackrel{!}{=} 0 \Rightarrow h(x) \equiv 0 \forall x \in \mathbb{R} \Rightarrow g(x) \equiv 0 \forall x \in \mathbb{R}$
Laplace Transform $\Rightarrow g(X) = 0$ with prob 1.

Side = German Tanks during WWII:

How many tanks are produced? $\rightarrow 1, 2, 3, \dots, N$: estimate N !

Y_n is sufficient for $N \Rightarrow \boxed{\hat{N} = \frac{n+1}{n} Y_n - 1}$ UMVUE for N
discrete case!

Terminology of "Completeness"

Def. We say that a statistic T is ancillary for θ if the distribution of T does not depend on θ .

Def. We say that a statistic T is 1st order ancillary for θ if $E(T)$ does not depend on θ .

It might seem that an ancillary statistic is useless when making inference about θ .

However, the reason for the term ancillary (servant, helper, assistant) is conveyed in the following example:

Ex: Toss a fair coin once: if the outcome is Heads (with prob $\frac{1}{2}$) then draw a sample of size n from a $N(\mu, 1)$ distribution. If the outcome is Tails: draw a sample of size n from a $N(\mu, \frac{1}{2})$ distribution.

$X :=$ r.v. defining the outcome of the coin toss:

$$\text{So } P(X=1) = \frac{1}{2} = P(X=0)$$

If we are interested in μ , then the statistic $T = X$ is ancillary for μ .
($\because P(X=x)$ does NOT depend on μ $\forall x \in \{0,1\}$)

However, surely the information that we get from X , does enhance our inference about μ .

By conditioning on the observed values of X , we gain information about the distribution from which our 2nd sample is drawn.

thus $\text{Var}(\bar{y}) = \sigma^2$ will change according to whether $X=1$ or $X=0$

A statistic T is 1st order ancillary $\Rightarrow E(T) \stackrel{\theta}{=} c \stackrel{\text{CST indep of } \theta}{\Rightarrow} E(T-c) \stackrel{\theta}{=} 0$

➔ completeness says that a family is complete if the only ancillary statistics are the trivial ones!

Method of Estimation:

Method of Moments:

→ Given a parameter θ to estimate: How do we derive a suitable estimator?

I The Method:

Setup: $X_1, \dots, X_n$ IID r.v. whose distribution depends on $\theta_1, \theta_2, \dots, \theta_r$.

- ALL k -moments exist for $1 \leq k \leq r$: i.e. $E[X^k] < \infty$.
- Let $\mu_k = E(X^k)$, $1 \leq k \leq r$, be the k^{th} population moment.
- Let $m_k := \frac{1}{n} \sum_{i=1}^n X_i^k$, $1 \leq k \leq r$, be the k^{th} sample moment.

WLLN: $\underbrace{\frac{1}{n} \sum_{i=1}^n X_i^k}_{m_k} \xrightarrow{n \rightarrow \infty} \underbrace{E[X^k]}_{\mu_k} \Rightarrow m_k \xrightarrow{n \rightarrow \infty} \mu_k$

Method: Set $m_k \approx \mu_k$, and solve: $\begin{cases} m_1 = h_1(\theta_1, \dots, \theta_r) \\ \vdots \\ m_r = h_r(\theta_1, \dots, \theta_r) \end{cases}$ for $\hat{\theta}_1, \dots, \hat{\theta}_r$.
(with $m_k = \frac{1}{n} \sum_{i=1}^n X_i^k$)

Result: You obtain: $\begin{cases} \hat{\theta}_1 = g_1(m_1, \dots, m_r) \\ \vdots \\ \hat{\theta}_r = g_r(m_1, \dots, m_r) \end{cases}$ } useful when g_k 's are continuous

Note: • The MOM is easy to apply (MOM estimators are often used as starting values in iterative methods)
• MOM's are not always unbiased.
• MOM's are consistent estimators! $\hat{\theta}_k \xrightarrow[n \rightarrow \infty]{P} \theta_k, \forall 1 \leq k \leq r$
EX: MLE's

(*) When g_k are continuous:

$$m_k = \frac{1}{n} \sum_{i=1}^n X_i^k \xrightarrow[n \rightarrow \infty]{P} E[X^k] = \mu_k \xRightarrow{g \text{ continuous}} \underbrace{g_k(m_1, \dots, m_r)}_{:= \hat{\theta}_k} \xrightarrow[n \rightarrow \infty]{P} g_k(\mu_1, \dots, \mu_r) = \theta_k \quad \forall 1 \leq k \leq r$$

II Examples:

Ex 1: $X_1, \dots, X_n$ IID $N(\mu, \sigma^2)$, $\mu \neq \sigma^2$ unknown

Solution: $\hat{\mu} = \frac{1}{n} \sum_{i=1}^n X_i$, $\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \rightarrow$ **BIASED!**

(*) $\theta = (\theta_1, \theta_2) = (\mu, \sigma^2)$ unknown: $\begin{cases} m_1 = \mu_1 \\ m_2 = \mu_2 \end{cases} \Rightarrow \begin{cases} \frac{1}{n} \sum_{i=1}^n X_i = \mu_1 = \mu \\ \frac{1}{n} \sum_{i=1}^n X_i^2 = \mu_2 = \sigma^2 + \mu^2 \end{cases}$
So: $\begin{cases} \mu = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X} \\ \sigma^2 = \left(\frac{1}{n} \sum_{i=1}^n X_i^2 \right) - \bar{X}^2 = \frac{1}{n} \left[\left(\sum_{i=1}^n X_i^2 \right) - n\bar{X}^2 \right] = \frac{1}{n} \sum_{i=1}^n (X_i^2 - \bar{X}^2) = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \end{cases}$

Ex 2: $X_1, \dots, X_n$ IID Bern(p), p unknown:

Solution: $\hat{p} = \frac{1}{n} \sum_{i=1}^n X_i$, where $\bar{X}$ = trial # of the 1st success

(*) $X \sim \text{Geom}(p) \Rightarrow m_1 = \mu_1 = E(X) = \frac{1}{p} \Rightarrow \hat{p} = \frac{1}{m_1} = \frac{1}{\bar{X}}$

OR: $\hat{p} = \frac{1}{n} \sum_{i=1}^n X_i$

(*) $P(X_i = 0) = 1-p$, $P(X_i = 1) = p$

So: $\begin{cases} m_1 = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X} \\ \mu_1 = E[X_i] = p \end{cases} \Rightarrow \hat{p} = \bar{X}$

Note: $\frac{1}{\bar{X}}$ a reasonable unbiased estimator for p :

Suppose $T(\bar{X})$ is an estimator for p :

Unbiasedness $\Rightarrow E[T(\bar{X})] = \sum_{x=1}^{\infty} T(x) p(1-p)^{x-1} \stackrel{!}{=} p \quad \forall p \Rightarrow \begin{cases} T(1) = 1 \\ T(x) = 0 \quad \forall x \geq 2 \end{cases} \Rightarrow$ crazy estimator!

Method of Maximum Likelihood :

① The Method :

Likelihood: • Discrete = $l(\theta; \tilde{x}) = P_{\tilde{x}}(\tilde{x}; \theta)$
 • Continuous = $l(\theta; \tilde{x}) = f_{\tilde{x}}(\tilde{x}; \theta)$ } likelihood of θ at the observation $\tilde{x}$
 $\Rightarrow l = \text{a function of } \underline{\theta}$
 (but l is a different function at each $\tilde{x}$)

MLE: $\hat{\theta} = \hat{\theta}(\tilde{x})$ is a MLE of θ if:

$$\hat{\theta}(\tilde{x}) = \arg \sup_{\theta \in (\Theta)} l(\theta; \tilde{x}), \quad \forall \tilde{x} \text{ s.t. } f_{\tilde{x}}(\tilde{x}; \theta) > 0$$

Prop: $\hat{\theta} = \hat{\theta}(\tilde{x})$ is a MLE $\Leftrightarrow \sup_{\theta \in (\Theta)} l(\theta; \tilde{x}) = l(\hat{\theta}(\tilde{x}), \tilde{x}) \quad \forall \tilde{x} \text{ s.t. } f_{\tilde{x}}(\tilde{x}; \theta) > 0$

Note = • the MLE may not exist

• the MLE is not always unique (but it often is!)

• $\hat{\theta}(\tilde{x})$ is the maximum likelihood estimator; $\hat{\theta}(\tilde{x})$ is the maximum likelihood estimate

Trick: To find $\hat{\theta}$, maximize $\log(l(\theta; \tilde{x}))$ WRT θ .

Method: • Continuous case = If l is differentiable, use calculus:

→ Solve the "likelihood equations": $\frac{\partial}{\partial \theta_i} \log(l(\tilde{\theta}; \tilde{x})) \stackrel{!}{=} 0$, for $\tilde{\theta} = (\theta_1, \dots, \theta_r)$
 → "Score vector": $(\frac{\partial}{\partial \theta_1} \log(l(\tilde{\theta}; \tilde{x})), \dots, \frac{\partial}{\partial \theta_r} \log(l(\tilde{\theta}; \tilde{x})))$

• discrete case = harder!

Ex 1: $X_1, \dots, X_n$ IID $\mathcal{N}(\mu, \sigma^2)$ r.v. with μ, σ^2 unknown $\Rightarrow$

$$\begin{cases} \hat{\mu} = \bar{X} \\ \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \end{cases}$$

$$(\otimes) \log(l(\mu, \sigma^2; \tilde{x})) = \log\left(\frac{1}{(2\pi)^{n/2}} \cdot \frac{1}{(\sigma^2)^{n/2}} e^{-\frac{1}{2\sigma^2} \sum_{i=1}^n (x_i - \mu)^2}\right) = \text{CST} - \frac{n}{2} \log(\sigma^2) - \frac{1}{2\sigma^2} \sum_{i=1}^n (x_i - \mu)^2$$

$$\bullet \frac{\partial}{\partial \mu} \log(l) \stackrel{!}{=} 0 \Rightarrow \sum_{i=1}^n (x_i - \mu) = 0 \Rightarrow \mu = \frac{1}{n} \sum_{i=1}^n x_i \quad \forall x_i \Rightarrow \hat{\mu} = \frac{1}{n} \sum_{i=1}^n x_i = \bar{X}$$

$$\bullet \frac{\partial}{\partial \sigma^2} \log(l) \stackrel{!}{=} 0 \Rightarrow -\frac{n}{2\sigma^2} + \frac{1}{2(\sigma^2)^2} \sum_{i=1}^n (x_i - \mu)^2 = 0 \Rightarrow \sigma^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2 \Rightarrow \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \hat{\mu})^2 = \bar{S}$$

Ex 2: $X_1, \dots, X_n$ IID Poisson(λ) r.v. with λ unknown:

$$\hat{\lambda} = \bar{X}$$

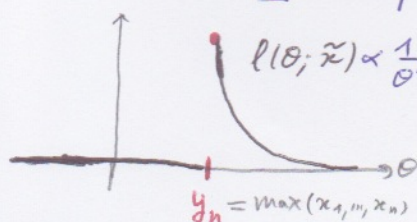
$$(\otimes) \log(l(\lambda; \tilde{x})) = \log\left(\frac{\lambda^{\sum_{i=1}^n x_i} e^{-n\lambda}}{\prod_{i=1}^n x_i!}\right) = \left(\sum_{i=1}^n x_i\right) \log \lambda - n\lambda - \text{CST}$$

$$\Rightarrow \frac{\partial}{\partial \lambda} \log(l) \stackrel{!}{=} 0 \Rightarrow \frac{1}{\lambda} \sum_{i=1}^n x_i - n = 0 \Rightarrow \lambda = \frac{1}{n} \sum_{i=1}^n x_i \quad \forall x_i \Rightarrow \hat{\lambda} = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X}$$

Ex 3: $X_1, \dots, X_n$ IID $U(0, \theta)$ r.v. with θ unknown:

$$\hat{\theta} = Y_n$$

→ Cannot set up the likelihood equations: l is NOT differentiable at the region of interest



Invariance Property of MLEs

Thm: Let θ an unknown parameter with MLE $\hat{\theta}$.
The MLE of $\eta = \tau(\theta)$ is: $\hat{\eta} = \tau(\hat{\theta})$
↳ some f of θ

$$\Rightarrow \boxed{\widehat{\tau(\theta)} = \tau(\hat{\theta})}$$

Ex: • $X_1, \dots, X_n \sim \text{Exp}(\beta)$: let $\theta = P[X > x] = e^{-x/\beta}$
 $\hat{\beta} = \bar{x} \Rightarrow \hat{\theta} = e^{-x/\hat{\beta}} \Rightarrow \widehat{P[X > x]} = e^{-x/\bar{x}}$

• $X_1, \dots, X_n \sim N(\mu, \sigma^2)$: let $\theta = \frac{\mu^2 + \sin \mu}{\log(\sigma^2)} \Rightarrow \hat{\theta} = \frac{\hat{\mu}^2 + \sin(\hat{\mu})}{\log(\hat{\sigma}^2)}$

Proof: Case 1: $\tau(\theta)$ is 1-1 $\sim \tau^{-1}$ exists!

Let $l^*(\eta; \tilde{x}) :=$ likelihood of η at $\tilde{x}$. and τ is 1-1 $\Rightarrow \tau(\theta) = \eta \Leftrightarrow \tau^{-1}(\eta) = \theta$

so $l^*(\eta; \tilde{x}) = l^*(\tau(\theta); \tilde{x}) = l(\theta; \tilde{x}) = l(\tau^{-1}(\eta); \tilde{x})$

$\Rightarrow l^*(\hat{\eta}; \tilde{x}) := \sup_{\eta} [l^*(\eta; \tilde{x})] = \sup_{\theta} [l(\theta; \tilde{x})] = l(\hat{\theta}; \tilde{x}) = l(\tau^{-1}(\tau(\hat{\theta})); \tilde{x}) = l^*(\tau(\hat{\theta}); \tilde{x})$

$\Rightarrow \hat{\eta} = \tau(\hat{\theta}) \Rightarrow \widehat{\tau(\theta)} = \tau(\hat{\theta})$

Case 2: $\tau(\theta)$ is not 1-1:

Induced likelihood: $l^*(\eta; \tilde{x}) := \sup_{\theta: \tau(\theta) = \eta} l(\theta; \tilde{x})$

so $l^*(\hat{\eta}; \tilde{x}) := \sup_{\eta} l^*(\eta; \tilde{x}) = \sup_{\eta} \sup_{\theta: \tau(\theta) = \eta} l(\theta; \tilde{x})$

iterated Supremum

$= \sup_{\theta} l(\theta; \tilde{x}) = l(\hat{\theta}; \tilde{x}) = \sup_{\theta: \tau(\theta) = \tau(\hat{\theta})} l(\theta; \tilde{x}) = l^*(\tau(\hat{\theta}); \tilde{x})$

$\Rightarrow l^*(\hat{\eta}; \tilde{x}) = l^*(\tau(\hat{\theta}); \tilde{x}) \Rightarrow \widehat{\tau(\theta)} = \tau(\hat{\theta})$

Efficiency of MLEs

Asymptotic Unbiasedness: $\hat{\theta}$ asympt. unbiased for $\theta \Leftrightarrow \exists K_n \in \mathbb{R}$ s.t. $K_n \cdot (\hat{\theta}_n - \theta) \xrightarrow[n \rightarrow \infty]{D} W \Rightarrow E(W) = 0$

! This says something about the limiting distrib of $\hat{\theta}_n$ ($\neq$ unbiased as $n \rightarrow \infty$: $E(\hat{\theta}_n) \xrightarrow[n \rightarrow \infty]{} \theta$)

Thm: $X_1, \dots, X_n$ IID r.v. with unknown Scalar parameter θ

$\hat{\theta}_n$ = MLE of θ_0 for each n :

then: $\sqrt{n} (\hat{\theta}_n - \theta) \xrightarrow[n \rightarrow \infty]{D} W \sim N(0, \frac{1}{I(\theta)})$, $I(\theta) = \text{Var} \left[\frac{\partial}{\partial \theta} \log(\theta, \tilde{x}) \right] \Big|_{\theta = \theta_0}$
↳ generic $\tilde{x}$
↳ CR Lower bound

corr: $\hat{\theta}_n$ is asymptotically unbiased

Note = MLEs are asymptotically efficient: $\text{Var}(\hat{\theta}_n) \simeq \frac{1}{n \cdot I(\theta_0)}$

$\rightarrow$ CR Lower bound!

MLEs have the lowest large Sample Variance than any other estimator!

Proof: $\begin{cases} L(\theta; \tilde{x}) := \log(l(\theta; \tilde{x})) \\ L'(\theta; \tilde{x}) := \frac{\partial}{\partial \theta} \log(\dots) \\ L''(\theta; \tilde{x}) := \frac{\partial^2}{\partial \theta^2} \log(\dots) \end{cases} \quad \hat{\theta}_n \xrightarrow[n \rightarrow \infty]{P} \theta \Rightarrow \underbrace{L'(\hat{\theta}_n; \tilde{x})}_{=0 \text{ as } \hat{\theta}_n \text{ is MLE}} \simeq L'(\theta; \tilde{x}) + (\hat{\theta}_n - \theta) L''(\theta; \tilde{x})$

so $L'(\theta; \tilde{x}) = -(\hat{\theta}_n - \theta) \cdot L''(\theta; \tilde{x}) \Rightarrow \sum_{i=1}^n L'(\theta; X_i) = -(\hat{\theta}_n - \theta) \cdot \sum_{i=1}^n L''(\theta; \tilde{x})$ ($\because \tilde{x} \sim \Pi \theta \Rightarrow l \rightarrow \Pi l$; $L \rightarrow \sum L_i$)

And: $\begin{cases} E_{\theta} [L'(\theta; X_i)] = \frac{\partial}{\partial \theta} \int \dots dx = 0 \quad \forall i \\ \text{Var}_{\theta} [L'(\theta; X_i)] = \text{Var}_{\theta} \left[\frac{\partial}{\partial \theta} \log(\theta, X_i) \right] \quad \forall i \end{cases} \Rightarrow \frac{\sum_{i=1}^n L'(\theta; X_i) - E[L'(\theta; X_i)]}{\sqrt{\text{Var}_{\theta} \left(\sum_{i=1}^n L'(\theta; X_i) \right)}} = \frac{-(\hat{\theta}_n - \theta) \cdot \sum_{i=1}^n L''(\theta; \tilde{x}) - 0}{\sqrt{n \cdot \text{Var}_{\theta} [L'(\theta; X_i)]}} \xrightarrow[n \rightarrow \infty]{D} N(0, 1)$

Also: $-\frac{1}{n} \sum_{i=1}^n L''(\theta; X_i) \xrightarrow[n \rightarrow \infty]{P} -E_{\theta} [L''(\theta; X_i)] = +E_{\theta} [L'(\theta; X_i)^2] = \text{Var}_{\theta} [L'(\theta; X_i)]$

$\Rightarrow \sqrt{n} (\hat{\theta}_n - \theta) \cdot \sqrt{\text{Var}_{\theta} [L'(\theta; X_i)]} \xrightarrow[n \rightarrow \infty]{D} N(0, 1) \Rightarrow \sqrt{n} (\hat{\theta}_n - \theta) \xrightarrow[n \rightarrow \infty]{D} N(0, \frac{1}{\text{Var} [L'(\theta; X_i)]})$

(multiply top & bottom by $\frac{1}{n}$)

Consistency of MLEs :

Consistency: $X_1, \dots, X_n$ r.v. with parameter θ :

An Estimator $\hat{\theta}_n$ is consistent for θ if: $\hat{\theta}_n \xrightarrow[n \rightarrow \infty]{P} \theta \quad \forall \theta \in \Theta$

Ex: $X_1, \dots, X_n \sim \mathcal{N}(\mu, \sigma^2)$ with (μ, σ^2) unknown: $\begin{cases} \hat{\mu}_n = \frac{1}{n} \sum_{i=1}^n X_i \\ \hat{\sigma}_n^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \end{cases}$ are consistent

(88) WLLN: $\hat{\mu}_n \xrightarrow[n \rightarrow \infty]{P} \mu$ and $\hat{\sigma}_n^2 \xrightarrow[n \rightarrow \infty]{P} \sigma^2$

Assumptions - To prove that $\hat{\theta}_n$ is consistent when $X_1, \dots, X_n$ are IID and Θ is finite, Need:

- A1: The family of distributions $\{P_\theta : \theta \in \Theta\}$ is identifiable: $\theta \neq \theta' \Rightarrow P_\theta \neq P_{\theta'}$
- A2: The Support of P_θ does not depend on θ .
- A3: X_i 's are IID
- A4: Θ Contains an open interval ST θ_0 is an interior point of this interval.

Lemma = $X_1, \dots, X_n$ IID r.v. with likelihood $l(\theta; \tilde{X})$; Let θ_0 be the "true" value of θ .

then: $P_{\theta_0} \left(l(\theta_0; \tilde{X}) > l(\theta; \tilde{X}) \right) \xrightarrow[n \rightarrow \infty]{} 1 \quad \forall \theta \in \Theta$

Note = Jensen's inequality: $\begin{cases} E[g(X)] \geq g(E(X)) & \text{if } g \text{ is convex} \\ E[g(X)] \leq g(E(X)) & \text{if } g \text{ is concave} \end{cases}$
 $\hookrightarrow$ Strict if Strictly convex/concave

Proof = $P_{\theta_0} (l(\theta_0; \tilde{X}) > l(\theta; \tilde{X})) \stackrel{\text{indep}}{=} P_{\theta_0} \left(\prod_{i=1}^n l(\theta_0; X_i) > \prod_{i=1}^n l(\theta; X_i) \right) = P_{\theta_0} \left(\prod_{i=1}^n \frac{l(\theta_0; X_i)}{l(\theta; X_i)} > 1 \right)$

So $P_{\theta_0} (l(\theta_0; \tilde{X}) > l(\theta; \tilde{X})) = P_{\theta_0} \left(\frac{1}{n} \sum_{i=1}^n \log \left(\frac{l(\theta_0; X_i)}{l(\theta; X_i)} \right) > 0 \right)$

• Let $Y_n := \frac{1}{n} \sum_{i=1}^n \log \left(\frac{l(\theta_0; X_i)}{l(\theta; X_i)} \right) \Rightarrow$ WLLN: $Y_n \xrightarrow[n \rightarrow \infty]{P} E_{\theta_0} \left[\log \left(\frac{l(\theta_0; X_1)}{l(\theta; X_1)} \right) \right] \quad (*)$

• Log is strictly concave $\Rightarrow$ Jensen: $E_{\theta_0} \left[\log \left(\frac{l(\theta_0; X_1)}{l(\theta; X_1)} \right) \right] < \log \left[E_{\theta_0} \left(\frac{l(\theta_0; X_1)}{l(\theta; X_1)} \right) \right] = \log \left[\int_{-\infty}^{\infty} \frac{l(\theta_0(x))}{l(\theta; x)} l(x; \theta_0) dx \right]$
 $= \log \left[\int_{-\infty}^{\infty} l(\theta_0(x)) dx \right] = \log[1] = 0$

So: $E_{\theta_0} \left[\log \left(\frac{l(\theta_0; X_1)}{l(\theta; X_1)} \right) \right] < 0, \quad \forall \theta \neq \theta_0$

• So $\forall \epsilon > 0$: $P_{\theta_0} [|Y_n - E_{\theta_0}[\dots]| \leq \epsilon] \xrightarrow[n \rightarrow \infty]{} 1$ by $(*)$

$\Rightarrow P_{\theta_0} [Y_n - E_{\theta_0}[\dots] \leq \epsilon] \xrightarrow[n \rightarrow \infty]{} 1$

$\Rightarrow P_{\theta_0} [Y_n \leq \epsilon + \underbrace{E_{\theta_0}[\dots]}_{< 0}] \xrightarrow[n \rightarrow \infty]{} 1 \quad \forall \epsilon > 0$

Pick ϵ small enough ST $\epsilon + E_{\theta_0}[\dots] < 0$

then: $P_{\theta_0} [Y_n < 0] \xrightarrow[n \rightarrow \infty]{} 1$

So $P_{\theta_0} [l(\theta_0; \tilde{X}) > l(\theta; \tilde{X})] \xrightarrow[n \rightarrow \infty]{} 1 \quad \forall \theta \in \Theta$

Thm : Let $\Theta = \{\theta_1, \dots, \theta_n\}$ be a finite discrete parameter space.

Then under $A_1, \dots, A_4$: (1) the MLE $\hat{\theta}_n$ for θ is consistent $\Leftrightarrow P_\theta [\hat{\theta}_n = \theta] \xrightarrow{n \rightarrow \infty} 1, \forall \theta \in \Theta$
 (2) The MLE $\hat{\theta}_n$ is consistent for θ : $\hat{\theta}_n \xrightarrow{P} \theta, \forall \theta \in \Theta$

Proof : (1) $P_\theta (|\hat{\theta}_n - \theta| < \epsilon) \xrightarrow{n \rightarrow \infty} 1 \quad \forall \epsilon > 0 \Leftrightarrow P_\theta (\hat{\theta}_n = \theta) \xrightarrow{n \rightarrow \infty} 1$

$\Rightarrow \Theta = \{\theta_1, \dots, \theta_n\} : wlog \theta_1 < \theta_2 < \dots < \theta_n$. Since $\theta \in \Theta$, $\exists i$ st $\theta = \theta_i \rightarrow$ pick $\epsilon < \theta_{i+1} - \theta_{i-1}$

$$\text{so } P_\theta (\theta - \epsilon < \hat{\theta}_n < \theta + \epsilon) = P_\theta (|\hat{\theta}_n - \theta| < \epsilon) \xrightarrow{n \rightarrow \infty} 1$$

$$= P_\theta (\hat{\theta}_n = \theta)$$

$$\text{so } \forall \epsilon > 0, \exists \epsilon_n > 0 \text{ st } P_\theta (\hat{\theta}_n = \theta) = P_\theta (|\hat{\theta}_n - \theta| < \epsilon_n) \Rightarrow P_\theta (\hat{\theta}_n = \theta) \xrightarrow{n \rightarrow \infty} 1$$

$$\Leftrightarrow \{\hat{\theta}_n : |\hat{\theta}_n - \theta| < \epsilon\} \supset \{\hat{\theta}_n : \hat{\theta}_n = \theta\} \Rightarrow 1 \geq P_\theta (|\hat{\theta}_n - \theta| < \epsilon) \geq P_\theta (\hat{\theta}_n = \theta) \xrightarrow{n \rightarrow \infty} 1$$

$$\rightarrow P_\theta (|\hat{\theta}_n - \theta| < \epsilon) \xrightarrow{n \rightarrow \infty} 1 \text{ by Squeeze Thm.}$$

(2) Particular case: $\Theta = \{\theta_1, \theta_2\}$

Enough to show that $P_{\theta_1} [\hat{\theta}_n = \theta_1] \xrightarrow{n \rightarrow \infty} 1$ and $P [\hat{\theta}_n = \theta_2] \xrightarrow{n \rightarrow \infty} 1$.

$$P_{\theta_1} [\hat{\theta}_n = \theta_1] = P_{\theta_1} [\hat{\theta}_n = \theta_1 | \ell(\theta_1; \tilde{x}) > \ell(\theta_2; \tilde{x})] \cdot P[\ell(\theta_1; \tilde{x}) > \ell(\theta_2; \tilde{x})] \\ + P_{\theta_1} [\hat{\theta}_n = \theta_1 | \ell(\theta_1; \tilde{x}) < \ell(\theta_2; \tilde{x})] \cdot P[\ell(\theta_1; \tilde{x}) < \ell(\theta_2; \tilde{x})]$$

$$\text{But: } P_{\theta_1} [\hat{\theta}_n = \theta_1 | \ell(\theta_1; \tilde{x}) > \ell(\theta_2; \tilde{x})] = P_{\theta_1} [\hat{\theta}_n = \theta_1 | \hat{\theta}_n = \theta_1] = 1$$

$\hookrightarrow \hat{\theta}_n = \theta_1 \text{ as MLE}$

$$P_{\theta_1} [\hat{\theta}_n = \theta_1 | \ell(\theta_1; \tilde{x}) < \ell(\theta_2; \tilde{x})] = P_{\theta_1} [\hat{\theta}_n = \theta_1 | \hat{\theta}_n = \theta_2] = 0 \text{ as } \theta_1 \neq \theta_2$$

$\hookrightarrow \hat{\theta}_n = \theta_2 \text{ as MLE}$

$$P_{\theta_1} [\ell(\theta_1; \tilde{x}) > \ell(\theta_2; \tilde{x})] \xrightarrow{n \rightarrow \infty} 1 \text{ by the previous lemma.}$$

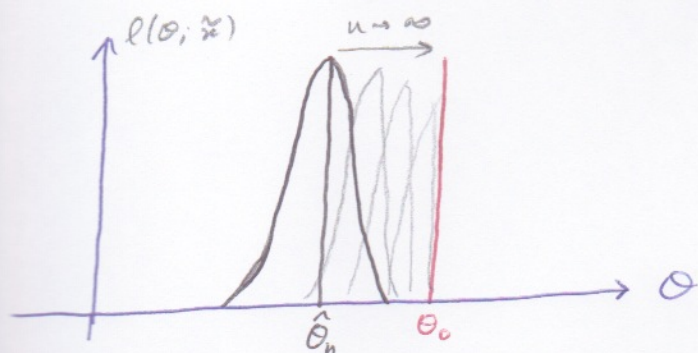
$$\text{so } P_{\theta_1} [\hat{\theta}_n = \theta_1] = 1 \cdot P_{\theta_1} [\ell(\theta_1; \tilde{x}) > \ell(\theta_2; \tilde{x})] + 0 \xrightarrow{n \rightarrow \infty} 1$$

(similar for $P_{\theta_2} [\hat{\theta}_n = \theta_2]$)

□

Thm : If $\ell(\theta; \tilde{x})$ is differentiable, then:

with probability 1, $\exists$ a consistent solution to the likelihood equations $\frac{\partial \log(\ell(\theta))}{\partial \theta} = 0$



Interval Estimation :

→ Confidence Intervals!

① General Picture :

$\tilde{X} = (x_1, \dots, x_n)$ a random vector, θ an unknown parameter :

- $L = L(\tilde{X})$ and $R = R(\tilde{X})$: left & Right random endpoints of the interval $(L(\tilde{X}), R(\tilde{X}))$
- $(L(\tilde{X}), R(\tilde{X}))$: random interval.

Goal = Find L & R , s.t. : given α , $P_\theta(L < \theta < R) = 1 - \alpha \quad \forall \theta \in \Theta$

Note : $(L(\tilde{X}), R(\tilde{X}))$ is really a probability interval : $\theta \in (L(\tilde{X}), R(\tilde{X}))$ with probability $(1 - \alpha)$!

• After the experiment, $\tilde{x} = (x_1, \dots, x_n)$ have fixed values, so $(L(\tilde{x}), R(\tilde{x}))$ is a fixed interval.
 $\hookrightarrow \theta \in (L(\tilde{x}), R(\tilde{x})) = 0 \text{ or } 1$!

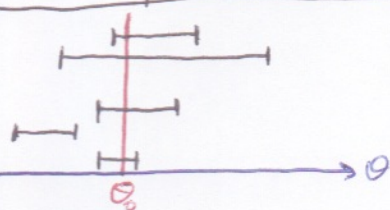
Coverage probability = $(1 - \alpha)$ → usually : $(1 - \alpha) \sim 0.95$
 (confidence level / coeff)

100(1 - α)% Confidence interval for $\theta = (L(\tilde{x}), R(\tilde{x}))$ as well as $(L(\tilde{X}), R(\tilde{X}))$
 → FIXED ! → RANDOM !

Weak interpretation of C.I. : → Repeat the experiment for the same θ

If we repeat the experiment many times with the same unknown θ , and compute the interval $(L(\tilde{x}), R(\tilde{x}))$ each time, then:

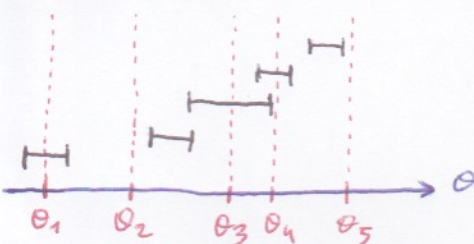
"100(1 - α)% of the times, roughly, we expect $\theta \in (L(\tilde{x}), R(\tilde{x}))$ ".



Strong interpretation of C.I. : → Repeat the experiment for possibly different θ 's.

If we use the random interval in a large nb of experiments, with possibly different θ 's, then:

"100(1 - α)% of the times, roughly, we expect $(L(\tilde{x}), R(\tilde{x}))$ to capture its respective θ ".



Reason: Let $Y_j = \begin{cases} 1, & \text{if } \theta_j \in I_j \\ 0, & \text{if } \theta_j \notin I_j \end{cases}$

where $I_j = j^{\text{th}}$ confidence interval.

• The j trials are independent $\Rightarrow Y_j$ are IID Bernoulli(p) r.v.

• $P = P_\theta[\theta_j \in I_j] = 1 - \alpha \quad \forall \theta_j$

• So $E[Y_j] = P = 1 - \alpha \Rightarrow \text{WLLN} : \frac{1}{n} \sum_{j=1}^n Y_j \xrightarrow[n \rightarrow \infty]{P} E[Y_j] = 1 - \alpha$

Proportion of Intervals that capture their θ_j 's.

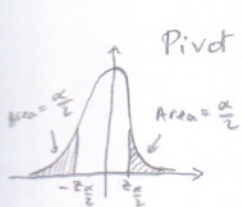
⚠ L & R should NOT depend on unknown parameters!

II The method of pivots:

Pivot = a random variable of the form $g(\tilde{X}; \theta)$ whose distribution does not depend on θ .
For each $\tilde{x}$, $g(\tilde{x}, \theta)$ is monotone in θ .

Setting 1: $X_1, \dots, X_n$ IID $\mathcal{N}(\mu, \sigma^2)$ r.v., σ^2 known $\Rightarrow \begin{cases} L = \bar{X} - z_{\alpha/2} \sigma/\sqrt{n} \\ R = \bar{X} + z_{\alpha/2} \sigma/\sqrt{n} \end{cases}$ for μ

Solution: Goal: find $L \& R$ st $P_\mu(L < \mu < R) = 1 - \alpha \quad \forall \mu$



Pivot: $\frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \sim \mathcal{N}(0,1)$ as $X_i \sim \mathcal{N}(\mu, \sigma^2)$, so: $P(-z_{\alpha/2} < Z < z_{\alpha/2}) = 1 - \alpha$
where $Z \sim \mathcal{N}(0,1)$

$$\text{So: } P_\mu\left(\frac{\bar{X} - R}{\sigma/\sqrt{n}} < \frac{\bar{X} - \mu}{\sigma/\sqrt{n}} < \frac{\bar{X} - L}{\sigma/\sqrt{n}}\right) = 1 - \alpha \Rightarrow \begin{cases} \frac{\bar{X} - R}{\sigma/\sqrt{n}} = -z_{\alpha/2} \\ \frac{\bar{X} - L}{\sigma/\sqrt{n}} = z_{\alpha/2} \end{cases} \Rightarrow \begin{cases} R = \bar{X} + z_{\alpha/2} \cdot \frac{\sigma}{\sqrt{n}} \\ L = \bar{X} - z_{\alpha/2} \cdot \frac{\sigma}{\sqrt{n}} \end{cases}$$

Setting 2: $X_1, \dots, X_n$ IID $\mathcal{N}(\mu, \sigma^2)$ r.v., σ^2 unknown $\Rightarrow \begin{cases} L = \bar{X} - t_{n-1}(\frac{\alpha}{2}) \cdot s/\sqrt{n} \\ R = \bar{X} + t_{n-1}(\frac{\alpha}{2}) \cdot s/\sqrt{n} \end{cases}$ for μ

Solution: Goal: find $L \& R$ st $P_\mu(L < \mu < R) = 1 - \alpha \quad \forall \mu$

Pivot: $\frac{\bar{X} - \mu}{s/\sqrt{n}} \sim t_{n-1} \Rightarrow P(-t_{n-1}(\frac{\alpha}{2}) < W < t_{n-1}(\frac{\alpha}{2})) = 1 - \alpha, W \sim t_{n-1}$

$$\text{So } P_\mu\left(\frac{\bar{X} - R}{s/\sqrt{n}} < \frac{\bar{X} - \mu}{s/\sqrt{n}} < \frac{\bar{X} - L}{s/\sqrt{n}}\right) = 1 - \alpha \Rightarrow \begin{cases} \frac{\bar{X} - R}{s/\sqrt{n}} = -t_{n-1}(\alpha/2) \\ \frac{\bar{X} - L}{s/\sqrt{n}} = t_{n-1}(\alpha/2) \end{cases}$$

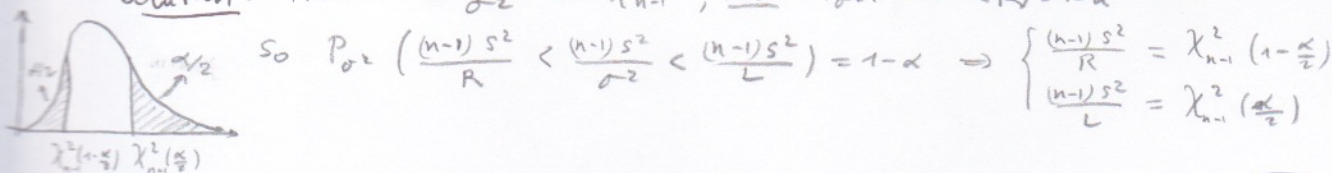
Setting 3: $X_1, \dots, X_n$ IID, $\mu \& \sigma^2$ unknown, $n \geq 30 \Rightarrow \begin{cases} L \approx \bar{X} - z_{\alpha/2} \cdot s/\sqrt{n} \\ R \approx \bar{X} + z_{\alpha/2} \cdot s/\sqrt{n} \end{cases}$ for μ

Solution: Pivot: $\frac{\bar{X} - \mu}{s/\sqrt{n}} \xrightarrow[n \rightarrow \infty]{\text{CLT}} \mathcal{N}(0,1)$

$$\text{So } P_\mu\left(\frac{\bar{X} - R}{s/\sqrt{n}} < \frac{\bar{X} - \mu}{s/\sqrt{n}} < \frac{\bar{X} - L}{s/\sqrt{n}}\right) = 1 - \alpha \Rightarrow \begin{cases} \frac{\bar{X} - R}{s/\sqrt{n}} \approx -z_{\alpha/2} \\ \frac{\bar{X} - L}{s/\sqrt{n}} \approx +z_{\alpha/2} \end{cases}$$

Setting 4: $X_1, \dots, X_n$ IID $\mathcal{N}(\mu, \sigma^2)$ r.v., $\mu \& \sigma^2$ unknown $\Rightarrow \begin{cases} L = s^2 \cdot \frac{n-1}{\chi_{n-1}^2(\alpha/2)} \\ R = s^2 \cdot \frac{n-1}{\chi_{n-1}^2(1-\alpha/2)} \end{cases}$ for σ^2

Solution: Pivot: $\frac{(n-1)s^2}{\sigma^2} \sim \chi_{n-1}^2$, Goal: $P_{\sigma^2}(L < \sigma^2 < R) = 1 - \alpha$



$$\text{So } P_{\sigma^2}\left(\frac{(n-1)s^2}{R} < \frac{(n-1)s^2}{\sigma^2} < \frac{(n-1)s^2}{L}\right) = 1 - \alpha \Rightarrow \begin{cases} \frac{(n-1)s^2}{R} = \chi_{n-1}^2(1-\frac{\alpha}{2}) \\ \frac{(n-1)s^2}{L} = \chi_{n-1}^2(\frac{\alpha}{2}) \end{cases}$$

Setting 5: $X_1, \dots, X_n$ IID Bernoulli(p) r.v. $\Rightarrow \begin{cases} L \approx \hat{p} - z_{\alpha/2} \sqrt{\hat{p}(1-\hat{p})/n} \\ R \approx \hat{p} + z_{\alpha/2} \sqrt{\hat{p}(1-\hat{p})/n} \end{cases}$ for p

Solution: Pivot: $\frac{\hat{p} - p}{\sqrt{\hat{p}(1-\hat{p})/n}} \xrightarrow[n \rightarrow \infty]{\text{CLT}} \mathcal{N}(0,1)$: Goal: $P_p(L < p < R) = 1 - \alpha$

$$\hookrightarrow \hat{p} := \frac{1}{n} \sum_{i=1}^n X_i$$

$$\text{So: } P_p\left(\frac{\hat{p} - R}{\sqrt{\hat{p}(1-\hat{p})/n}} < \frac{\hat{p} - p}{\sqrt{\hat{p}(1-\hat{p})/n}} < \frac{\hat{p} - L}{\sqrt{\hat{p}(1-\hat{p})/n}}\right) = 1 - \alpha \Rightarrow \begin{cases} \frac{\hat{p} - R}{\sqrt{\hat{p}(1-\hat{p})/n}} \approx -z_{\alpha/2} \\ \frac{\hat{p} - L}{\sqrt{\hat{p}(1-\hat{p})/n}} \approx +z_{\alpha/2} \end{cases}$$

Setting 6: $X_1, \dots, X_n$ IID r.v., μ known $\Rightarrow \begin{cases} L \approx \\ R \approx \end{cases}$

Solution: Pivot: $\frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \xrightarrow[n \rightarrow \infty]{\text{CLT}} \mathcal{N}(0,1)$: Goal: $P_{\sigma^2}(L < \sigma^2 < R) = 1 - \alpha$

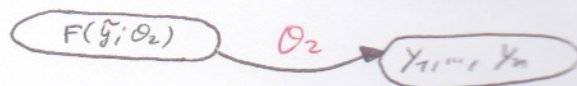
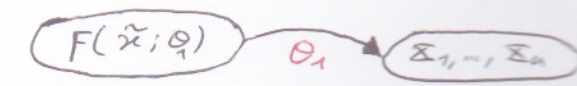
$$\text{So } P_{\sigma^2}\left(\frac{\bar{X} - \mu}{\sqrt{R/n}} < \frac{\bar{X} - \mu}{\sigma/\sqrt{n}} < \frac{\bar{X} - \mu}{\sqrt{L/n}}\right) = 1 - \alpha \Rightarrow \begin{cases} \frac{\bar{X} - \mu}{\sqrt{R/n}} \approx -z_{\alpha/2} \\ \frac{\bar{X} - \mu}{\sqrt{L/n}} \approx +z_{\alpha/2} \end{cases} \Rightarrow \begin{cases} R \approx \frac{(\bar{X} - \mu)^2 n}{(z_{\alpha/2})^2} \\ L \approx \frac{(\bar{X} - \mu)^2 n}{(z_{\alpha/2})^2} \end{cases} \dots ?$$

Require $\bar{X} > \mu$

III 2-Samples Confidence Intervals :

Idea : • up to now : "single-sample problems" → draw a sample from a single population & make inference about it

• "2-samples problems" → compare 2 populations : make inference about $(\mu_1 - \mu_2)$, $\frac{\sigma_1^2}{\sigma_2^2}$, $(p_1 - p_2)$, $\frac{p_1}{p_2}$



Compare θ_1 & θ_2 !

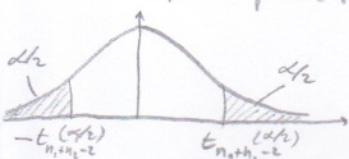
Setting 1 : $\tilde{x}_1, \dots, \tilde{x}_{n_1}$ IID $\mathcal{P}(\mu_1, \sigma^2)$
 $\tilde{y}_1, \dots, \tilde{y}_{n_2}$ IID $\mathcal{P}(\mu_2, \sigma^2)$
 $\tilde{x} \perp \tilde{y}$, (μ_1, μ_2, σ^2) unknown

$\Rightarrow P_{\mu_1, \mu_2} (L < \mu_1 - \mu_2 < R) = 1 - \alpha$, $\forall \mu_1 \in \Theta_1$
 $\forall \mu_2 \in \Theta_2$

unbiased : $E[S_p^2] = \frac{(n_1-1)\sigma^2 + (n_2-1)\sigma^2}{n_1+n_2-2} = \sigma^2$

• Pivot : $U = \frac{(\bar{x} - \bar{y}) - (\mu_1 - \mu_2)}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}$ where $S_p^2 = \frac{(n_1-1)S_1^2 + (n_2-1)S_2^2}{n_1+n_2-2}$ → "Pooled Sample Variance"

If we prove that $U \sim t_{n_1+n_2-2}$: $P_{\mu_1, \mu_2} \left(\frac{(\bar{x} - \bar{y}) - R}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} < \frac{(\bar{x} - \bar{y}) - (\mu_1 - \mu_2)}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} < \frac{(\bar{x} - \bar{y}) - L}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} \right) = 1 - \alpha$



$\Rightarrow \begin{cases} \frac{(\bar{x} - \bar{y}) - R}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} = -t_{n_1+n_2-2}(\alpha/2) \\ \frac{(\bar{x} - \bar{y}) - L}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} = t_{n_1+n_2-2}(\alpha/2) \end{cases}$

$\Rightarrow \begin{cases} L = (\bar{x} - \bar{y}) - t_{n_1+n_2-2}(\alpha/2) \cdot S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}} \\ R = (\bar{x} - \bar{y}) + t_{n_1+n_2-2}(\alpha/2) \cdot S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}} \end{cases}$

• Claim : $\frac{(\bar{x} - \bar{y}) - (\mu_1 - \mu_2)}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} \sim t_{n_1+n_2-2}$

Goal : write $U \sim \frac{\mathcal{N}(0,1)}{\sqrt{\chi_{n_1+n_2-2}^2/(n_1+n_2-2)}}$ Independent RV.

Step 1 : $U = \frac{(\bar{x} - \bar{y}) - (\mu_1 - \mu_2)}{\sigma \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} \bigg/ \sqrt{\frac{(n_1+n_2-2)S_p^2}{\sigma^2} / (n_1+n_2-2)} = \frac{V}{W}$

Step 2 : $V \sim \mathcal{N}(0,1)$ ($\because \tilde{x} \perp \tilde{y} \Rightarrow \bar{x} \perp \bar{y} \Rightarrow (\bar{x} - \bar{y})$ is Normal as $\bar{x}$ & $\bar{y}$ are normal.

$E[\bar{x} - \bar{y}] = E[\bar{x}] - E[\bar{y}] = \mu_1 - \mu_2$

$Var[\bar{x} - \bar{y}] = Var[\bar{x}] + Var[\bar{y}] = \frac{\sigma^2}{n_1} + \frac{\sigma^2}{n_2} = \sigma^2 \left(\frac{1}{n_1} + \frac{1}{n_2} \right)$ } $V \sim \mathcal{N}(0,1)$

Step 3 : $W \sim \chi_{n_1+n_2-2}^2 / (n_1+n_2-2)$

(*) $\frac{(n_1+n_2-2)S_p^2}{\sigma^2} = \frac{(n_1-1)S_1^2}{\sigma^2} + \frac{(n_2-1)S_2^2}{\sigma^2} \sim \chi_{n_1-1}^2 + \chi_{n_2-1}^2 \sim \chi_{n_1+n_2-2}^2$ because $\tilde{x} \perp \tilde{y} \Rightarrow S_1^2 \perp S_2^2$
 So both χ^2 are indep.

Step 4 : $V \perp W$, i.e. $(\bar{x} - \bar{y}) \perp S_p^2 \rightarrow$ Goal : $f_{\bar{x}, \bar{y}, S_1^2, S_2^2}(r, s, u, v) = f_{\bar{x}, \bar{y}}(r, s) \cdot f_{S_1^2, S_2^2}(u, v)$ $\forall r, s, u, v$

$f_{\bar{x}, \bar{y}, S_1^2, S_2^2}(r, s, u, v) = f_{(\bar{x}, \bar{y}), (S_1^2, S_2^2)}(r, s, u, v) = f_{\bar{x}, S_1^2}(r, u) \cdot f_{\bar{y}, S_2^2}(s, v) = f_{\bar{x}}(r) f_{S_1^2}(u) f_{\bar{y}}(s) f_{S_2^2}(v) = f_{\bar{x}}(r) f_{\bar{y}}(s) f_{S_1^2}(u) f_{S_2^2}(v) = f_{\bar{x}, \bar{y}}(r, s) f_{S_1^2, S_2^2}(u, v)$
 $\uparrow$ $\uparrow$
 $\tilde{x} \perp \tilde{y}$ $\tilde{x} \perp S_1^2$ $\tilde{y} \perp S_2^2$ (under Normality) $\tilde{x} \perp \tilde{y}$

so $(\bar{x}, \bar{y}) \perp (S_1^2, S_2^2)$

$\Rightarrow (\bar{x} - \bar{y}) \perp S_p^2$ as funct^o of independent vars are indep.

Setting 2: $X_1, \dots, X_{n_1}$ IID r.v. with (μ_1, σ_1^2)
 $Y_1, \dots, Y_{n_2}$ IID r.v. with (μ_2, σ_2^2)
 $\bar{X} \perp \bar{Y}$ and $(\mu_1, \mu_2, \sigma_1^2, \sigma_2^2)$ unknown

$$\left. \begin{array}{l} X_1, \dots, X_{n_1} \text{ IID r.v. with } (\mu_1, \sigma_1^2) \\ Y_1, \dots, Y_{n_2} \text{ IID r.v. with } (\mu_2, \sigma_2^2) \\ \bar{X} \perp \bar{Y} \text{ and } (\mu_1, \mu_2, \sigma_1^2, \sigma_2^2) \text{ unknown} \end{array} \right\} n_1, n_2 \geq 30 \Rightarrow \begin{cases} L \approx (\bar{X} - \bar{Y}) - z_{\frac{\alpha}{2}} \sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}} \\ R \approx (\bar{X} - \bar{Y}) + z_{\frac{\alpha}{2}} \sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}} \end{cases}$$

for $P_{\mu_1, \mu_2} (L < \mu_1 - \mu_2 < R) = 1 - \alpha$

• Pivot: $Z = \frac{(\bar{X} - \bar{Y}) - (\mu_1 - \mu_2)}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}} = Z$

If we prove that $Z \xrightarrow[n \rightarrow \infty]{D} \mathcal{N}(0,1)$, then $P_{\mu_1, \mu_2} \left(\underbrace{\frac{(\bar{X} - \bar{Y}) - R}{\sqrt{\dots}}}_{= -z_{\frac{\alpha}{2}}} < \frac{(\bar{X} - \bar{Y}) - (\mu_1 - \mu_2)}{\sqrt{\dots}} < \underbrace{\frac{(\bar{X} - \bar{Y}) - L}{\sqrt{\dots}}}_{= +z_{\frac{\alpha}{2}}} \right) = 1 - \alpha$

• claim: $\frac{(\bar{X} - \bar{Y}) - (\mu_1 - \mu_2)}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}} \xrightarrow[n_1, n_2 \rightarrow \infty]{D} \mathcal{N}(0,1)$

(oo) $\frac{\bar{X} - \mu_1}{\sqrt{s_1^2/n_1}} \xrightarrow[n_1 \rightarrow \infty]{D} \mathcal{N}(0,1)$ and $\frac{\bar{Y} - \mu_2}{\sqrt{s_2^2/n_2}} \xrightarrow[n_2 \rightarrow \infty]{D} \mathcal{N}(0,1)$ by CLT

So $\frac{\bar{X} - \mu_1}{\sqrt{s_1^2/n_1}} \cdot \frac{\sqrt{s_1^2/n_1}}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}} + \frac{\bar{Y} - \mu_2}{\sqrt{s_2^2/n_2}} \cdot \frac{\sqrt{s_2^2/n_2}}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}} \xrightarrow[n_1, n_2 \rightarrow \infty]{D} \mathcal{N}(0,1)$ as $\bar{X} \perp \bar{Y}$ so indep normals.

$\frac{1}{\sqrt{1 + \frac{s_2^2/n_2}{s_1^2/n_1}}} \xrightarrow[n_1 \rightarrow \infty]{D} 1$ $\frac{1}{\sqrt{1 + \frac{s_1^2/n_1}{s_2^2/n_2}}} \xrightarrow[n_2 \rightarrow \infty]{D} 1$

Setting 3: $X_1, \dots, X_n$ IID Bernoulli(p_1)
 $Y_1, \dots, Y_n$ IID Bernoulli(p_2)
 $\bar{X} \perp \bar{Y}$ and (p_1, p_2) unknown

$$\left. \begin{array}{l} X_1, \dots, X_n \text{ IID Bernoulli}(p_1) \\ Y_1, \dots, Y_n \text{ IID Bernoulli}(p_2) \\ \bar{X} \perp \bar{Y} \text{ and } (p_1, p_2) \text{ unknown} \end{array} \right\} n_1, n_2 \geq 30 \Rightarrow \begin{cases} \hat{p}_1 = \frac{1}{n_1} \sum_{i=1}^{n_1} X_i \\ \hat{p}_2 = \frac{1}{n_2} \sum_{i=1}^{n_2} Y_i \end{cases}$$

$$\begin{cases} L \approx (\hat{p}_1 - \hat{p}_2) - z_{\frac{\alpha}{2}} \sqrt{\frac{\hat{p}_1(1-\hat{p}_1)}{n_1} + \frac{\hat{p}_2(1-\hat{p}_2)}{n_2}} \\ R \approx (\hat{p}_1 - \hat{p}_2) + z_{\frac{\alpha}{2}} \sqrt{\frac{\hat{p}_1(1-\hat{p}_1)}{n_1} + \frac{\hat{p}_2(1-\hat{p}_2)}{n_2}} \end{cases}$$

for $P_{p_1, p_2} (L < p_1 - p_2 < R) = 1 - \alpha$

• Pivot: $Z = \frac{(\hat{p}_1 - \hat{p}_2) - (p_1 - p_2)}{\sqrt{\frac{\hat{p}_1(1-\hat{p}_1)}{n_1} + \frac{\hat{p}_2(1-\hat{p}_2)}{n_2}}} \xrightarrow[n_1, n_2 \rightarrow \infty]{D} \mathcal{N}(0,1)$ by CLT

• So: $P_{p_1, p_2} \left(\underbrace{\frac{(\hat{p}_1 - \hat{p}_2) - R}{\sqrt{\dots}}}_{\approx -z_{\alpha/2}} < \frac{(\hat{p}_1 - \hat{p}_2) - (p_1 - p_2)}{\sqrt{\dots}} < \underbrace{\frac{(\hat{p}_1 - \hat{p}_2) - L}{\sqrt{\dots}}}_{\approx z_{\alpha/2}} \right) = 1 - \alpha$

(oo) $E[\hat{p}_1 - \hat{p}_2] = E\left[\frac{1}{n_1} \sum_{i=1}^{n_1} X_i\right] - E\left[\frac{1}{n_2} \sum_{i=1}^{n_2} Y_i\right] = p_1 - p_2$
 $\hat{p}_1 \perp \hat{p}_2$
 $\text{Var}[\hat{p}_1 - \hat{p}_2] = \text{Var}[\hat{p}_1] + \text{Var}[\hat{p}_2] = \frac{1}{n_1} p_1(1-p_1) + \frac{1}{n_2} p_2(1-p_2)$

Setting 4:

Example 7: Confidence interval for the ratio σ_1^2/σ_2^2 of variances of two normal distributions $N(\mu_1, \sigma_1^2)$ and $N(\mu_2, \sigma_2^2)$.

We shall assume we have independent random samples of sizes n_1 and n_2 from these distributions, and that μ_1 and μ_2 are unknown. From chapter 1, the quantities $Y_1 = \frac{(n_1-1)s_1^2}{\sigma_1^2}$ and $Y_2 = \frac{(n_2-1)s_2^2}{\sigma_2^2}$ have chi-square distributions with $n_1 - 1$ and $n_2 - 1$ d.f. respectively, and so

$$\frac{Y_2/(n_2 - 1)}{Y_1/(n_1 - 1)} = \frac{s_2^2/\sigma_2^2}{s_1^2/\sigma_1^2}$$

has the F -distribution with $n_2 - 1, n_1 - 1$ degrees of freedom. Hence our pivot will be

$$\frac{s_2^2/\sigma_2^2}{s_1^2/\sigma_1^2} \sim F_{n_2-1, n_1-1}.$$

which has the F -distribution with $n_2 - 1, n_1 - 1$ degrees of freedom. We have

$$\begin{aligned} 1 - \alpha &= P \left[F_{1-\alpha/2, n_2-1, n_1-1} < \frac{s_2^2/\sigma_2^2}{s_1^2/\sigma_1^2} < F_{\alpha/2, n_2-1, n_1-1} \right] \\ &= P \left[\frac{s_1^2}{s_2^2} \cdot F_{1-\alpha/2, n_2-1, n_1-1} < \frac{\sigma_1^2}{\sigma_2^2} < \frac{s_1^2}{s_2^2} \cdot F_{\alpha/2, n_2-1, n_1-1} \right] \end{aligned}$$

and so we find the confidence interval to be

$$\boxed{\frac{s_1^2}{s_2^2} \cdot F_{1-\alpha/2, n_2-1, n_1-1} < \frac{\sigma_1^2}{\sigma_2^2} < \frac{s_1^2}{s_2^2} \cdot F_{\alpha/2, n_2-1, n_1-1}.$$

Sometimes the fact that

$$F_{1-\alpha/2, n_2-1, n_1-1} = \frac{1}{F_{\alpha/2, n_1-1, n_2-1}} \quad \text{is used in this c.i.}$$

Population(s)	To test	calculate	and reject H_0 at level α if	$(1 - \alpha) \times 100\%$ Confidence Interval
1. $N(\mu, \sigma^2)$ (or any population if n is large (≥ 25)) with σ^2 known.	$H_0: \mu = \mu_0$ $H_1: \mu \neq \mu_0$ ($\mu > \mu_0$) ($\mu < \mu_0$)	$z = \frac{\bar{x} - \mu_0}{\sigma/\sqrt{n}}$ If n is large, then σ may be replaced by s .	$ z > z_{\alpha/2}$ ($z > z_{\alpha}$) ($z < -z_{\alpha}$)	$\bar{x} - z_{\alpha/2} \frac{\sigma}{\sqrt{n}} < \mu < \bar{x} + z_{\alpha/2} \frac{\sigma}{\sqrt{n}}$ If n is large, then σ may be replaced by s .
2. $N(\mu, \sigma^2)$ with both μ and σ^2 unknown	same as above	$t = \frac{\bar{x} - \mu_0}{s/\sqrt{n}}$	$ t > t_{\alpha/2, n-1}$ ($t > t_{\alpha, n-1}$) ($t < -t_{\alpha, n-1}$)	$\bar{x} - t_{\alpha/2, n-1} \frac{s}{\sqrt{n}} < \mu < \bar{x} + t_{\alpha/2, n-1} \frac{s}{\sqrt{n}}$
3. Bin(p) (large sample case ($n \geq 25$)). $\hat{p} = x/n, \hat{q} = 1 - \hat{p}$	$H_0: p = p_0$ $H_1: p \neq p_0$ ($p > p_0$) ($p < p_0$)	$z = \frac{\hat{p} - p_0}{\sqrt{\frac{p_0 q_0}{n}}}$	$ z > z_{\alpha/2}$ ($z > z_{\alpha}$) ($z < -z_{\alpha}$)	$\hat{p} - z_{\alpha/2} \sqrt{\frac{\hat{p}\hat{q}}{n}} < p < \hat{p} + z_{\alpha/2} \sqrt{\frac{\hat{p}\hat{q}}{n}}$
4. $N(\mu, \sigma^2)$ with both μ and σ^2 unknown	$H_0: \sigma^2 = \sigma_0^2$ $H_1: \sigma^2 \neq \sigma_0^2$ ($\sigma^2 > \sigma_0^2$) ($\sigma^2 < \sigma_0^2$)	$\chi^2 = \frac{(n-1)s^2}{\sigma_0^2}$	$\chi^2 > \chi_{\alpha/2, n-1}^2$ or $\chi^2 < \chi_{1-\alpha/2, n-1}^2$ ($\chi^2 > \chi_{\alpha, n-1}^2$) ($\chi^2 < \chi_{1-\alpha, n-1}^2$)	$\frac{(n-1)s^2}{\chi_{\alpha/2, n-1}^2} < \sigma^2 < \frac{(n-1)s^2}{\chi_{1-\alpha/2, n-1}^2}$
5. $N(\mu_1, \sigma_1^2)$ and $N(\mu_2, \sigma_2^2)$ (or any populations provided n_1 and n_2 are large (≥ 25)), where σ_1^2 and σ_2^2 are known.	$H_0: \mu_1 - \mu_2 = D_0$ $H_1: \mu_1 - \mu_2 \neq D_0$ ($\mu_1 - \mu_2 > D_0$) ($\mu_1 - \mu_2 < D_0$)	$z = \frac{\bar{x}_1 - \bar{x}_2 - D_0}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}}$ If n_1 and n_2 are large, σ_1^2 and σ_2^2 may be replaced by s_1^2 and s_2^2	$ z > z_{\alpha/2}$ ($z > z_{\alpha}$) ($z < -z_{\alpha}$)	$\bar{x}_1 - \bar{x}_2 - z_{\alpha/2} \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}} < \mu_1 - \mu_2 < \bar{x}_1 - \bar{x}_2 + z_{\alpha/2} \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}$ If n_1 and n_2 are large, σ_1^2 and σ_2^2 may be replaced by s_1^2 and s_2^2
6. $N(\mu_1, \sigma_1^2)$ and $N(\mu_2, \sigma_2^2)$, where σ_1^2 and σ_2^2 are unknown, but we assume $\sigma_1^2 = \sigma_2^2$.	same as above	$t = \frac{\bar{x}_1 - \bar{x}_2 - D_0}{s_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}$ where $s_p^2 = \frac{(n_1-1)s_1^2 + (n_2-1)s_2^2}{n_1 + n_2 - 2}$	$ t > t_{\alpha/2, n_1+n_2-2}$ ($t > t_{\alpha, n_1+n_2-2}$) ($t < -t_{\alpha, n_1+n_2-2}$)	$\bar{x}_1 - \bar{x}_2 - t_{\alpha/2, n_1+n_2-2} s_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}} < \mu_1 - \mu_2 < \bar{x}_1 - \bar{x}_2 + t_{\alpha/2, n_1+n_2-2} s_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}$
7. Bin(p_1) and Bin(p_2) (large sample case-both $n_1, n_2 \geq 25$) $\hat{p}_1 = x_1/n_1, \hat{q}_1 = 1 - \hat{p}_1$ $\hat{p}_2 = x_2/n_2, \hat{q}_2 = 1 - \hat{p}_2$	$H_0: p_1 - p_2 = D_0$ $H_1: p_1 - p_2 \neq D_0$ ($p_1 - p_2 > D_0$) ($p_1 - p_2 < D_0$)	$z = \begin{cases} \frac{\hat{p}_1 - \hat{p}_2 - D_0}{\sqrt{\frac{\hat{p}_1 \hat{q}_1}{n_1} + \frac{\hat{p}_2 \hat{q}_2}{n_2}}} & \text{if } D_0 \neq 0, \\ \frac{\hat{p}_1 - \hat{p}_2}{\sqrt{\frac{\hat{p} \hat{q}}{n_1} + \frac{\hat{p} \hat{q}}{n_2}}} & \text{if } D_0 = 0, \end{cases}$ where $\hat{p} = \frac{x_1 + x_2}{n_1 + n_2}$	$ z > z_{\alpha/2}$ ($z > z_{\alpha}$) ($z < -z_{\alpha}$)	$\hat{p}_1 - \hat{p}_2 - z_{\alpha/2} \sqrt{\frac{\hat{p}_1 \hat{q}_1}{n_1} + \frac{\hat{p}_2 \hat{q}_2}{n_2}} < p_1 - p_2 < \hat{p}_1 - \hat{p}_2 + z_{\alpha/2} \sqrt{\frac{\hat{p}_1 \hat{q}_1}{n_1} + \frac{\hat{p}_2 \hat{q}_2}{n_2}}$
8. $N(\mu_1, \sigma_1^2)$ and $N(\mu_2, \sigma_2^2)$, where μ_1 and μ_2 are unknown.	$H_0: \sigma_1^2 = \sigma_2^2$ $H_1: \sigma_1^2 \neq \sigma_2^2$ ($\sigma_1^2 > \sigma_2^2$) ($\sigma_1^2 < \sigma_2^2$)	$F = \frac{s_1^2}{s_2^2}$ Take pop'n 1 to be that with the larger sample variance.	$F \geq F_{\alpha/2, n_1-1, n_2-1}$ ($F > F_{\alpha, n_1-1, n_2-1}$) do not reject	$\frac{s_1^2}{s_2^2} \frac{1}{F_{\alpha/2, n_1-1, n_2-1}} < \frac{\sigma_1^2}{\sigma_2^2} < \frac{s_1^2}{s_2^2} F_{\alpha/2, n_2-1, n_1-1}$

If samples are assumed randomly drawn. In the two-population scenarios 5,6,7, and 8, the two samples are assumed independent of each other.

2.1.2 Method of Moments.

This method is based on the strong law of large numbers: if $X_1, X_2, X_3, \dots$ is a sequence of i.i.d. random variables with finite k th moment $\mu'_k = E(X^k)$, then

$$\frac{1}{n} \sum_{i=1}^n X_i^k \rightarrow \mu'_k \quad \text{as } n \rightarrow \infty.$$

Thus if $X_1, \dots, X_n$ is a random sample and n is large, and if we put

$$m_k = \frac{1}{n} \sum_{i=1}^n X_i^k,$$

then we have

$$m_k \approx \mu'_k, \quad k \geq 1.$$

The method of moments consists of solving as many of the equations

$$m_k = \mu'_k, \quad k \geq 1,$$

starting with the case $k = 1$, as are necessary to identify the unknown parameters.

Example. Suppose $X_1, \dots, X_n$ is a random sample from an exponential distribution with mean θ . Then $\mu'_1 = \theta$ and $m_1 = \bar{X}$, so the moment estimator is $\hat{\theta} = \bar{X}$.

Example. Find the moment estimator of θ given a random sample $X_1, \dots, X_n$ from the Bernoulli distribution

$$P[X = 0] = 1 - \theta, \quad P[X = 1] = \theta.$$

Solution. we have $m_1 = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X}$ and $\mu'_1 = EX = \theta$, so the moment estimator of θ is $\bar{X}$.

Example. Suppose $X_1, \dots, X_n$ is a random sample from a gamma distribution with parameters α and β . Find moment estimators of α and β .

Solution. We have $\mu'_1 = \alpha\beta$ and $\mu'_2 = E(X^2) = \text{Var}(X) + (EX)^2 = \alpha\beta^2 + \alpha^2\beta^2$. Hence $m_1 = \alpha\beta$ and $m_2 = \alpha\beta^2(1 + \alpha)$. Solving these two equations for α and β , we find

$$\hat{\alpha} = \frac{m_1^2}{m_2 - m_1^2}, \quad \hat{\beta} = \frac{m_2 - m_1^2}{m_1}.$$