

Laurent DELSOL

laurent.delsol@univ-orleans.fr.

19. 10. 2022.

Estimates methods in Statistics.

Part 1. Parametric estimation. $\theta \in \mathbb{R}^d$.

Definition

Property.

Part 2. Non parametric estimation.

f_X r.

Exercises

Practice on R software.

8h - 11h15

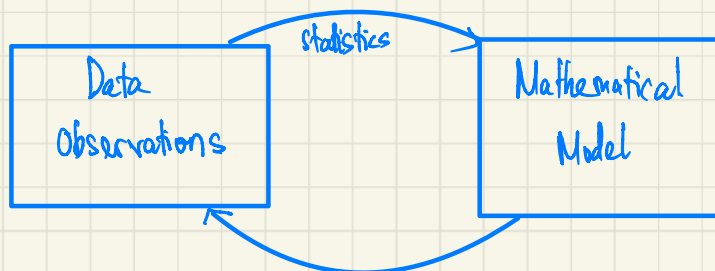
13h - 16h15.

Part 1

Parametric Estimation.

Introduction. Probability vs Statistics.

Example. Flip a coin. $\begin{cases} 1: \text{tail} \\ 0: \text{else.} \end{cases}$



$\left\{ \begin{array}{l} \text{Probabilistic model } (\Omega, \mathcal{F}, \mathbb{P}) \\ \text{Statistical model } (\Omega, \mathcal{F}, \{\mathbb{P}_\theta, \theta \in \Theta\}) \end{array} \right.$

kg chứa tham số.

Estimation problem. Data get something about θ .

Statistical framework:

One observes n random variables $X_1, X_2, \dots, X_n$ taking value in a measurable space $(X, \mathcal{F})$. We will assume in this lecture that those variables are independent with the same law $\mathbb{P}_{X, \theta}$. We will write here after

$X_1, X_2, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} \mathbb{P}_{X, \theta}$. $(X_1, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} X)$. And $\mathbb{P}_{X, \theta}$ is unknown.

And X would represent a variable whose distribution is $\mathbb{P}_{X, \theta}$.

i.i.d: independent and identically distributed.

Aim: Is it possible to infer something on $\mathbb{P}_{X, \theta}$ or to get from the observation of the random variable $X_1, \dots, X_n$?

$X = X_1, \dots, X_n$: random sample.

X_1 : represent the behavior of the first individual.

X_j : $\xrightarrow{j \text{th}}$

Observation: $x_1, \dots, x_n$ values taken by $X_1, \dots, X_n$ in the study.

A. Definitions.

Def 1.1. The sequences of RVs $(X_1, \dots, X_n)$ is a random sample of size n if $X_1, \dots, X_n \stackrel{i.i.d}{\sim} P_X$

. The distribution of the sample is $(P_X)^{\otimes n}$

. When the statistical study is performed you observe the values of these RVs $x_1, \dots, x_n$ (DATA)

Def 1.2. A statistical model is said to be **parametric** if the set of parameters $\Theta \subseteq \mathbb{R}^m$ ($m < +\infty$)
 $(\Omega, \mathcal{A}, \{P_\theta, \theta \in \Theta \subseteq \mathbb{R}^m\})$

parameter
is vector

↓
biết phân phối gì nhưng có
tham số chưa biết.

Def 1.3. A parametric model is said to be **dominated** if there exists a σ -finite dominant measure Q in such a way that all probabilities P_θ are absolutely continuous with respect to Q .

. We will then write $h(\cdot, \theta)$ the density of P_θ with respect to Q and hence the density of the random sample $\underline{X} = (X_1, \dots, X_n)$ is:

$$h(\underline{x}, \theta) = h(x_1, \theta) \dots h(x_n, \theta).$$

Examples.

continuous RV $\rightarrow Q$: Lebesgue measure.

discrete RV $\rightarrow Q$: counting measure $\mathbb{N}$.

We will denote hereafter E_θ and Var_θ the expectation and variance to explain they refer to P_θ .

Examples Discrete case.

Flip a coin $P_\theta = \text{Ber}(\theta)$ $X \sim \text{Ber}(\theta)$.

$P(X=1) = \theta$ X takes value in $\{0,1\}$.

$$P(X=0) = 1-\theta.$$

$Q = \sum_{k \in \mathbb{N}} \delta_k$: δ_k is Dirac distribution at k .
counting measure.

$$Qf = \sum_{k \in \mathbb{N}} \delta_k f = \sum_{k \in \mathbb{N}} f(k).$$

If $P_\theta = \text{Ber}(\theta)$ $(H) =]0,1[$.

$$P_{(H)} = \{ \text{Ber}(\theta), \theta \in (H) =]0,1[\}$$

$$h(x, \theta) = \theta^x (1-\theta)^{1-x} = \begin{cases} \theta & \text{if } x=1, \\ 1-\theta & \text{if } x=0. \end{cases}$$

$$\begin{aligned} h_n(x, \theta) &= \theta^{x_1} \dots \theta^{x_n} (1-\theta)^{1-x_1} \dots (1-\theta)^{1-x_n} \\ &= \theta^{y_n} (1-\theta)^{n-y_n}, \text{ where } y_n = x_1 + \dots + x_n. \end{aligned}$$

$$Y_n = \sum_{i=1}^n X_i \sim \text{Bin}(n, \theta) \quad \text{for binomial} \quad \binom{n}{x} \theta^x (1-\theta)^{n-x}.$$

Continuous case. $P_\theta = \mathcal{N}(\mu, \sigma)$, $\theta = (\mu, \sigma) \in \mathbb{R} \times \mathbb{R}_+^*$

σ : standard deviation

Q: Lebesgue measure.

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

$$h_n(x, \theta) = \left(\frac{1}{\sigma\sqrt{2\pi}} \right)^n \cdot e^{-\sum_{i=1}^n (x_i - \mu)^2 / 2\sigma^2}$$

Notation The function $\theta \mapsto h_n(x, \theta)$, when x is fixed is called the likelihood of the sample (usually written $L_n(x, \theta)$). Of course the likelihood depends on dominant measure.

Def 1.4. A statistic T is a function $\Omega \rightarrow \mathbb{R}^k$ computed from the random sample $X_1, \dots, X_n$

Ex $X_1, \dots, X_n$ random sample

$$T = \frac{X_1 + \dots + X_n}{n} = \bar{X}$$

$$w \mapsto (X_1(w), \dots, X_n(w))$$

$$\mapsto T(X_1(w), \dots, X_n(w))$$

$$f: \mathbb{R}^n \rightarrow \mathbb{R}^k$$

$$T(w) = f(X_1(w), \dots, X_n(w))$$

$$T = X_1, T = 1, T = \bar{X}, T = X, T = \min(X_i, 1 \leq i \leq n) \dots$$

1.2. Parametric estimators.

We consider a parametric statistical model $\mathcal{P} = \{P_\theta, \theta \in \Theta\}$, $\Theta \subseteq \mathbb{R}^m$ and a random sample $X_1, \dots, X_n \stackrel{i.i.d.}{\sim} P_\theta \in \mathcal{P}$.

Def. Estimator of $\lambda = g(\theta)$ (g is a known function). θ : natural parameter
 λ : interest parameter

An estimator of λ usually denoted $\hat{\lambda}$, is a statistic computed from the sample $X_1, \dots, X_n$ (and taking values in a $g(\Theta)$)

Remark. This definition is very general and does not ensure at all that $\hat{\lambda}$ will be closed to λ .

Ex Flip a coin $P_\theta = \text{Ber}(\theta)$, $\theta \in]0, 1[$

$$\lambda = \theta.$$

$$\hat{\lambda} = 1/4 \text{ or } \hat{\lambda} = X_1 \text{ or } \hat{\lambda} = \bar{X} \rightarrow \text{đều ok theo def, but}$$

$\downarrow$ a.s
 $E[X] = \theta = \lambda$ $\rightarrow$ gần $\lambda \rightarrow$ rất ok.

all test need in this case $\theta = \lambda$.

1.2.1 Expected properties

a) Bias

The bias of an estimator $\hat{\lambda}$ of λ is the quantity:

$$\text{Bias}_\theta(\hat{\lambda}) = E_\theta[\hat{\lambda}] - \lambda$$

trung bình của $\hat{\lambda}$. unbiased



biased.

An estimator is said to be unbiased if $\text{Bias}_\theta(\hat{\lambda}) = 0$

biased if $\text{Bias}_\theta(\hat{\lambda}) \neq 0$.

asymptotically unbiased if $\text{Bias}_\theta(\hat{\lambda}) \rightarrow 0$
 $n \rightarrow +\infty$

Example $X \sim \text{Ber}(\theta)$, $X_1, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} X$

$$\hat{\theta}_1 = X_1$$

$$\hat{\theta}_2 = \bar{X}$$

$$\hat{\theta}_3 = 1/4$$

Biased?

$$\cdot E_{\theta}[\hat{\theta}_1] = E_{\theta}[X_1] = \theta$$

$$\cdot E_{\theta}[\hat{\theta}_2] = E_{\theta}\left[\frac{1}{n} \sum_{i=1}^n X_i\right] = \frac{1}{n} \sum_{i=1}^n E_{\theta}[X_i] = \frac{1}{n} \cdot \sum_{i=1}^n E[X] = E[X] = \theta$$

↳ Note: $E_{\theta}[\bar{X}] = E_{\theta}[X]$ for all cases.

$$E_{\theta}[1/4] = 1/4 \neq \theta$$

→ $\hat{\theta}_1, \hat{\theta}_2$: unbiased estimator

$\hat{\theta}_3$: biased estimator.

Example. $X \sim \mathcal{N}(\mu, \sigma)$, $\lambda = \sigma^2$.

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

$$\hat{\sigma}^2 = \frac{1}{n} \cdot \sigma^2 \cdot Y_n$$

$$\rightarrow E_{\theta}[\hat{\sigma}^2] = \frac{1}{n} \cdot \sigma^2 \cdot E(Y_n) = \frac{n-1}{n} \sigma^2$$

$$E_{\theta}[\hat{\sigma}^2] = \frac{n-1}{n} \sigma^2$$

$$E(Y_n) = n-1$$

since $Y_n = \sum_{i=1}^n \left(\frac{X_i - \bar{X}}{\sigma}\right)^2 \sim \chi^2(n-1)$: χ^2 chi bình phương.

$$\cdot \text{Bias}(\hat{\sigma}^2) = \frac{n-1}{n} \sigma^2 - \sigma^2 = -\frac{1}{n} \sigma^2 \xrightarrow{n \rightarrow \infty} 0$$

↳ $\hat{\sigma}^2$ is an asymptotically unbiased.

$$\Rightarrow \hat{\sigma}_{n-1}^2 = \frac{n}{n-1} \hat{\sigma}^2 \quad : \text{d.f. unbiased} \rightarrow \hat{\sigma}_{n-1}^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

Then

$$E_{\theta}[\hat{\sigma}_{n-1}^2] = E_{\theta}\left[\frac{n}{n-1} \hat{\sigma}^2\right]$$

$$= \frac{n}{n-1} E_{\theta}[\hat{\sigma}^2]$$

$$= \frac{n}{n-1} \cdot \frac{n-1}{n} \sigma^2 = \sigma^2 \rightarrow \text{unbiased.}$$

b) Compare variances.

If $\hat{\lambda}_1$ and $\hat{\lambda}_2$ are 2 unbiased estimators of λ . Which is the best?



small variance
Better



great variance
Worse.

. Quadratic risk.

$$E_{\theta}[(\hat{\lambda} - \lambda)^2] = \text{Var}_{\theta}(\hat{\lambda}) + (\text{Bias}_{\theta}(\hat{\lambda}))^2$$

. For unbiased estimator, the quadratic risk is equal to the variance of $\hat{\lambda}$

• Hence the best estimator will be the one with smallest variance.

Def. If $\lambda = g(\theta) \in \mathbb{R}$

λ is square integrable.

Then the quadratic risk $R_\theta(\hat{\lambda})$ is $R_\theta(\hat{\lambda}) = E_\theta[(\hat{\lambda} - \lambda)^2]$

Property

$$R_\theta(\hat{\lambda}) = \text{Var}(\hat{\lambda}) + (\text{Bias}(\hat{\lambda}))^2.$$

$R_\theta(\hat{\lambda}) \xrightarrow{n \rightarrow \infty} 0$ then L^2 consistent.

New unbiased then $R_\theta(\hat{\lambda}) = \text{Var}(\hat{\lambda})$.

Come back to our example

$$X \sim \text{Ber}(\theta)$$

$$\begin{aligned} \hat{\theta}_1 &= X_1 \\ \hat{\theta}_2 &= \bar{X} \end{aligned} \quad \left. \vphantom{\begin{aligned} \hat{\theta}_1 &= X_1 \\ \hat{\theta}_2 &= \bar{X} \end{aligned}} \right\} \text{unbiased}$$

$$\text{Var}_\theta(\hat{\theta}_1) = \theta(1-\theta) \rightarrow \not\rightarrow 0 \Rightarrow \text{not } L^2 \text{ consistent.}$$

$$\begin{aligned} \text{Var}_\theta(\hat{\theta}_2) &= \text{Var}\left(\frac{1}{n} \sum_{i=1}^n X_i\right) \\ &\stackrel{\text{independent}}{=} \frac{1}{n^2} \sum_{i=1}^n \text{Var}(X_i) \end{aligned}$$

$$= \frac{1}{n^2} \sum_{i=1}^n \text{Var}(X) = \frac{\text{Var}(X)}{n} = \frac{\theta(1-\theta)}{n}.$$

Remarks: $X_1 \dots X_n \stackrel{\text{i.i.d.}}{\sim} X$

$$E_\theta[\bar{X}] = E_\theta[X]$$

$$\text{Var}_\theta(\bar{X}) = \frac{\text{Var}_\theta(X)}{n}$$

$$R(\hat{\theta}_1) = \theta(1-\theta)$$

$$R(\hat{\theta}_2) = \frac{\theta(1-\theta)}{n} \xrightarrow{n \rightarrow \infty} 0$$

$\hat{\theta}_2 \xrightarrow{L^2} \bar{X}$.
 $\hookrightarrow \hat{\theta}_2$ is better.

$\hat{\lambda}$ is said to be **consistent** in quadratic mean if $R(\hat{\lambda}) \xrightarrow[n \rightarrow \infty]{} 0$
(converge in L^2 sense)

c) Consistent estimator
Our aim $\hat{\lambda} \xrightarrow[n \rightarrow \infty]{} \lambda$?
Hầu kiểu gì cũng đc

- L^2 sense
- a.s
- P : hầu theo XS
- Q : P^2

Def. Let $\hat{\lambda}$ be (a sequence of) estimators of λ . Then $\hat{\lambda}$ is a **strongly consistent estimator** of λ if
 $\forall \epsilon \in \mathbb{R}$ $\hat{\lambda} \xrightarrow[n \rightarrow \infty]{P_0 \text{ a.s.}} \lambda$ dùng SLLN.

Remark. Strong Law of Large Numbers (SLLN)

If $\begin{cases} X_1, \dots, X_n \text{ i.i.d } X, X \in \mathbb{R}^k \\ \varphi: \mathbb{R}^k \rightarrow \mathbb{R}^n \end{cases}$

with $E[|\varphi(X)|] < +\infty$

$$\overline{\varphi(X)} = \frac{\varphi(X_1) + \dots + \varphi(X_n)}{n} \xrightarrow[n \rightarrow \infty]{P_0 \text{ a.s.}} E[|\varphi(X)|]$$

Come back to our example

$X \sim \text{Ber}(\theta)$

dùng $P_0(\hat{\lambda})$. (có đvnh n mẫu nui đến htn)

	Unbiased	L^2	consistent	P_0 a.s consistent
$\hat{\theta}_1 = X_1$	OK	No	No	No.
$\hat{\theta}_2 = \bar{X}$	OK	OK	OK	OK (dùng SLLN)
$\hat{\theta}_3 = 1/4$	No	No	No	No

If $\lambda = \theta(1-\theta) = g(\theta)$ then $\hat{\lambda} = \bar{X}(1-\bar{X})$ is strongly consistent
 because $\begin{cases} g \text{ is continuous} \\ \bar{X} \text{ is strongly consistent} \end{cases}$ from the continuous mapping theorem.

$$\left. \begin{array}{l} \bar{X} \text{ htv hkn v} \hat{\theta} \\ g \text{ true} \end{array} \right\} \Rightarrow g(\bar{X}) \xrightarrow{\mathbb{P}_{\theta} \text{ a.s.}} g(\theta).$$

d) Asymptotic distribution

Asymptotic normality.

Def Let $\hat{\lambda}$ be a sequence of estimations of λ . The sequence is said

asymptotically normally distributed if for all $\theta \in \Theta$

$$\begin{aligned} \sqrt{n}(\hat{\lambda} - \lambda) &\xrightarrow[n \rightarrow +\infty]{\mathcal{L}} \mathcal{N}(0, \Sigma(\theta)) \\ \text{or} \quad \sqrt{n}(\hat{\lambda} - \lambda) &\xrightarrow[n \rightarrow +\infty]{} \mathcal{N}(0, \Sigma(\theta)) \end{aligned}$$

↑ variance-covariance matrix.

Example. $X \sim \text{Ber}(\theta)$

$$\hat{\theta}_2 = \bar{X}, \quad \mathbb{E}_{\theta}[X] = \theta, \quad \text{Var}(X) = \theta(1-\theta)$$

$$\frac{\bar{X} - \theta}{\sqrt{\frac{\theta(1-\theta)}{n}}} \xrightarrow[n \rightarrow +\infty]{\mathcal{L}} \mathcal{N}(0, 1) \quad \text{Gaussian.} \quad \begin{array}{l} \text{(Central Limit Theorem)} \\ \text{CLT} \end{array}$$

$$\sqrt{n}(\bar{X} - \theta) \xrightarrow[n \rightarrow +\infty]{\mathcal{L}} \mathcal{N}(0, \theta(1-\theta))$$

Main idea. CLT + Δ method

Slutsky theorem.

1.3. Comparing estimators and optimization.

The aim is to compare estimators and get the best of them.

Def A loss function L is an application $L: \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}_+$ such that

$\forall x \in \mathbb{R}^d, L(x, x) = 0$ and / when in addition

$\forall x, y \in \mathbb{R}^d, x \neq y \Rightarrow L(x, y) > 0$ | we talk about a contrast function

Example. $L_p(x, y) = \sum_{i=1}^d |x_i - y_i|^p$

. If $p = 2$, the quadratic loss function. More generally L is convex and gives zero when $x = y$.

. If $\hat{\lambda}$ is an estimator of λ then $L(\hat{\lambda}, \lambda)$ is the loss coming from the estimator of λ .

Remark. It is usually difficult to compare directly $L(\hat{\lambda}, \lambda)$ and $L(\tilde{\lambda}, \lambda)$ because both of them are RVs.

. For instance if $\lambda \in \mathbb{R}$: $L_2(\hat{\lambda}, \lambda) = (\hat{\lambda} - \lambda)^2$

$$L_2(\tilde{\lambda}, \lambda) = (\tilde{\lambda} - \lambda)^2$$

Def. A risk is the expectation of a loss function

$$R(\hat{\lambda}) = \mathbb{E}_{\theta} [L(\hat{\lambda}, \lambda)]$$

Remarks

1) Quadratic risk is linked to the quadratic loss function.

$$2) R(\hat{\lambda}) = \text{Var}_\theta(\hat{\lambda}) + (\text{Bias}(\hat{\lambda}))^2$$

Def: 1) $\hat{\lambda}$ is better than $\tilde{\lambda}$ for R if and only if $\forall \theta \in \Theta: R(\hat{\lambda}) \leq R(\tilde{\lambda})$

2) strictly better <

Usually it is not easy to optimize a risk function.

In this lecture we mainly consider quadratic risk.

2. How to construct interesting estimators?

2.1. Unbiased linear estimator

Let $X_1, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} P_\theta, \theta \in \Theta$ is unknown.

Assume that for all $\theta \in \Theta$, the moment $m_\theta = \mathbb{E}_\theta[X]$ and $\mathbb{E}_\theta[X^2]$ are finite. $\sigma_\theta^2 = \text{Var}_\theta(X) = \mathbb{E}_\theta[X^2] - (\mathbb{E}_\theta[X])^2$

We consider linear estimator of $\lambda = m_\theta$:

$$\hat{\lambda} = b + \sum_{i=1}^n a_i X_i$$

Proposition For a n -sample with at least second order moment, the
arithmetic empirical mean $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$ is the best unbiased
estimator of m_θ for the quadratic risk. (if there exists $\theta_1, \theta_2: m_{\theta_1} \neq m_{\theta_2}$)

Proof To get an unbiased estimator, we want $\mathbb{E}_\theta[\hat{\lambda}] = \lambda$.

$$\begin{aligned}
 \text{Now } E_{\theta}[\hat{\lambda}] &= E_{\theta}\left[b + \sum_{i=1}^n a_i X_i\right] \\
 &= b + \sum_{i=1}^n a_i E_{\theta}[X_i] \\
 &= b + \sum_{i=1}^n a_i E_{\theta}[X] \\
 &= b + m_{\theta} \cdot \sum_{i=1}^n a_i
 \end{aligned}$$

For our estimator to be unbiased we want:

$$b + m_{\theta} \cdot \sum_{i=1}^n a_i = m_{\theta} \quad \text{for any } \theta \in \Theta$$

Now take θ_1, θ_2 such $m_{\theta_1} \neq m_{\theta_2}$, we have:

$$\begin{aligned}
 \begin{cases} m_{\theta_1} = b + m_{\theta_1} \cdot \sum_{i=1}^n a_i \\ m_{\theta_2} = b + m_{\theta_2} \cdot \sum_{i=1}^n a_i \end{cases} &\Leftrightarrow \begin{cases} m_{\theta_1} = b + m_{\theta_1} \cdot \sum_{i=1}^n a_i \\ m_{\theta_2} - m_{\theta_1} = (m_{\theta_2} - m_{\theta_1}) \cdot \sum_{i=1}^n a_i \end{cases} \\
 \Leftrightarrow \begin{cases} b = 0 \\ \sum_{i=1}^n a_i = 1 \end{cases}
 \end{aligned}$$

Then $\hat{\lambda} = \sum_{i=1}^n a_i X_i$ where $\sum_{i=1}^n a_i = 1$.

The quadratic risk is then:

$$\begin{aligned}
 E_{\theta}[(\hat{\lambda} - \lambda)^2] &= \text{Var}_{\theta}(\hat{\lambda}) + \underbrace{\left(\text{Bias}_{\theta}(\hat{\lambda})\right)^2}_{=0 \text{ (by unbiased)}} \\
 &= \text{Var}_{\theta}(\hat{\lambda})
 \end{aligned}$$

$$\text{Var}_{\theta}(\hat{\lambda}) = \text{Var}_{\theta}\left(\sum_{i=1}^n a_i X_i\right) = \sum_{i=1}^n a_i^2 \text{Var}_{\theta}(X) = \sum_{i=1}^n a_i^2 \sigma_{\theta}^2$$

Since

$$1 = 1^2 = \left(\sum a_i\right)^2 \stackrel{\text{(Cauchy-Schwarz)}}{\leq} \left(\sum a_i^2\right) \cdot \left(\sum 1^2\right) = n \cdot \sum a_i^2$$

$$\Rightarrow \sum a_i^2 \geq \frac{1}{n}$$

Then the quadratic risk is $R_2(\lambda) \geq \frac{1}{n} \cdot \sigma_0^2$

If we take $a_i = 1/n$ for $0 \leq i \leq n$ then $\sum a_i^2 = 1/n$

↳ Mean is the best unbiased estimator (linear) $\square$

Remark. In Cauchy Schwarz inequality, we have the equality iff

$$(a_1, a_2, \dots, a_n) = k \cdot (1, \dots, 1) \text{ for some value } k \in \mathbb{R},$$

which means

$$\forall i: 1 \leq i \leq n : a_i = k.$$

There is no other minimizer.

2.2. Method of moments.

- Starting point: Strong Law of Large Numbers

- We will use this with $\varphi: x \mapsto x^p$

if $E[|x|^p] < +\infty$

$$\overset{\substack{\text{theoretical} \\ p^{\text{th}} \text{ empirical} \\ \text{moment}}}{\overline{x^p}} = \frac{x_1^p + \dots + x_n^p}{n} \xrightarrow[n \rightarrow +\infty]{\text{a.s.}} E[x^p] =: \mu_p(\theta) \overset{p^{\text{th}} \text{ theoretical} \text{ moment}}{}$$

Algorithm

- 1) compute the first moment μ_1, μ_2 — in function of λ . báo nhiên từ đó cần để để tính bất biến μ_k
- 2) solve the equation to get λ in functions of $\mu_1, \mu_2, \dots$
- 3) Get $\hat{\lambda}$ by replacing μ_p by $\overline{x^p}$

Come back to examples

1) $X \sim \text{Ber}(\theta)$.

1) $\mu_1 = E_\theta[X] = \theta$.

2) $\theta = \mu_1 = E_\theta[X]$

3) $\hat{\theta} = \bar{X}$

đếm số lần thất bại từ khi thành công

2) $X \sim G(p)$ $P(X=k) = (1-p)^k \cdot p$

1) $\mu_1 = E_\theta[X] = \frac{1-p}{p} = \frac{1}{p} - 1$.

2) $p = \frac{1}{\mu_1 + 1} = \frac{1}{E_\theta[X] + 1}$

3) $\hat{p} = \frac{1}{\bar{X} + 1}$.

3) $X \sim \mathcal{N}(\mu, \sigma)$ $\lambda = (\mu, \sigma^2)$

1) $\begin{cases} \mu_1 = E_\theta[X] = \mu \\ \mu_2 = E_\theta[X^2] = \mu^2 + \sigma^2 \end{cases}$

2) $\begin{cases} \mu = \mu_1 \\ \sigma^2 = \mu_2 - \mu_1^2 \end{cases}$

3) $\begin{cases} \hat{\mu} = \bar{X} \\ \hat{\sigma}^2 = \bar{X^2} - (\bar{X})^2 \end{cases}$

$$4) X \sim \mathcal{E}(\theta) : p^{\text{mũ}} \quad f_X(x) = \theta \cdot e^{-\theta x} \mathbb{1}_{x \geq 0}.$$

$$1) \mu_1 = E_\theta[X] = \frac{1}{\theta}$$

$$2) \theta = \frac{1}{\mu_1}$$

$$3) \hat{\theta} = \frac{1}{\bar{X}}$$

↳ nhanh nhing = tĩt bĩng maximum likelihood.

2.3. Maximum likelihood

We assume to have a dominated model.

$$\mathbb{P}_\Theta = \{ \mathbb{P}_\theta, \theta \in \Theta \}.$$

in which for all $\theta \in \Theta$, $\mathbb{P}_\theta \ll \mathbb{Q}$. (with $\mathbb{Q}$: σ -finite positive measure)

Examples

$$1) \Theta = \{ \theta_1, \theta_2 \}.$$

$$\mathbb{Q} = \mathbb{P}_{\theta_1} + \mathbb{P}_{\theta_2}.$$

$$\text{then } \begin{cases} \mathbb{P}_{\theta_1} \ll \mathbb{Q} \\ \mathbb{P}_{\theta_2} \ll \mathbb{Q} \end{cases}$$

đĩc tĩyĩt đĩĩ: $\mathbb{Q}(A) = 0 \Rightarrow \mathbb{P}_\theta(A) = 0$

$$2) \mathcal{Q}_1(\mathbb{R}, \mathcal{B}(\mathbb{R})), \mathbb{P}_\theta = \mathcal{N}(0, 1), \theta \in \mathbb{R}.$$

$\mathbb{Q} = \lambda$ Lebesgue measure

$$d\mathbb{P}_\theta(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{(x-\theta)^2}{2}} d\lambda(x).$$

Def. The likelihood of a random sample $X_1, \dots, X_n$ is given by
 $\theta \mapsto L(\theta, X_1, \dots, X_n) = \prod_{i=1}^n h(x_i, \theta)$ häm merst ab xs f_x

Remark. For all θ , $L(\theta, X_1, \dots, X_n)$ is a RV.

Def. The Maximum likelihood estimator of λ is the value of λ maximizing
 $\lambda \mapsto L(\theta, X_1, \dots, X_n)$

$$\hat{\lambda}_{MLE} = \arg \max_{\lambda} L(\theta, X_1, \dots, X_n)$$

Remark.

- 1) $\hat{\lambda}$ may be non unique or may not exist.
- 2) In practice, when for all $\theta \in (\Theta)$, almost surely $L(\theta, X_1, \dots, X_n) > 0$, we take the logarithm:

$$\begin{aligned} \ln L(\theta, X_1, \dots, X_n) &= \ln \left(\prod_{i=1}^n h(x_i, \theta) \right) \\ &= \sum_{i=1}^n \ln(h(x_i, \theta)) \end{aligned}$$

Then we usually search λ such that: ($\lambda \in \mathbb{R}$)

$$\frac{\partial}{\partial \lambda} \ln L(\theta, X_1, \dots, X_n) = 0$$

$$\frac{\partial^2}{\partial \lambda^2} \ln L(\theta, X_1, \dots, X_n) < 0$$

20.10.2022.

Examples

1) $X \sim \text{Ber}(\theta)$, $\theta \in]0,1[$

$$h(x; \theta) = \theta^x (1-\theta)^{1-x}.$$

$$\begin{aligned} L(\theta, X_1, \dots, X_n) &= \prod_{i=1}^n \theta^{x_i} (1-\theta)^{1-x_i} \\ &= \theta^{\sum x_i} (1-\theta)^{n - \sum x_i} \\ &= \theta^{n\bar{x}} (1-\theta)^{n(1-\bar{x})} \end{aligned}$$

$$\begin{aligned} \Rightarrow LL(\theta, X_1, \dots, X_n) &= \ln(\theta^{n\bar{x}}) + \ln((1-\theta)^{n(1-\bar{x})}) \\ &= n\bar{x} \cdot \ln \theta + n(1-\bar{x}) \cdot \ln(1-\theta). \end{aligned}$$

$$\begin{aligned} \frac{\partial}{\partial \theta} LL(\theta, X_1, \dots, X_n) &= \frac{n\bar{x}}{\theta} - \frac{n(1-\bar{x})}{1-\theta} \\ &= \frac{n}{\theta(1-\theta)} [\bar{x}(1-\theta) - \theta(1-\bar{x})] \\ &= \frac{n(\bar{x} - \theta)}{\theta(1-\theta)} \end{aligned}$$

$$\frac{\partial}{\partial \theta} LL(\theta, X_1, \dots, X_n) = 0 \Leftrightarrow \bar{x} = \theta.$$

$$\frac{\partial^2}{\partial \theta^2} LL = \frac{-n\bar{x}}{\theta^2} - \frac{n(1-\bar{x})}{(1-\theta)^2} < 0$$

Hence $\hat{\theta}_{MLE} = \bar{X}$

⚠ MLE is better than moment estimator.
(thường ra kết quả giống nhau).

$$2) X \sim \mathcal{E}(\lambda) \quad h(x, \lambda) = \lambda \cdot e^{-\lambda x} \mathbb{1}_{x \geq 0}$$

$$\begin{aligned} L(\lambda, x_1, \dots, x_n) &= \prod_{i=1}^n h(x_i, \lambda) \\ &= \lambda^n \cdot e^{-\lambda \sum_{i=1}^n x_i} \prod_{i=1}^n \mathbb{1}_{x_i \geq 0} \\ &= \lambda^n \cdot e^{-n\lambda \bar{x}} \boxed{\prod_{i=1}^n \mathbb{1}_{x_i \geq 0}} \\ &= \lambda^n \cdot e^{-n\lambda \bar{x}} \cdot \mathbb{1}_{[\min x_i \geq 0]} \stackrel{as.}{=} 1 \end{aligned}$$

$$\cdot \mathcal{L}(\lambda, x_1, \dots, x_n) = n \cdot \ln \lambda - n\lambda \bar{x}$$

$$\frac{\partial}{\partial \lambda} \mathcal{L} = \frac{n}{\lambda} - n\bar{x}$$

$$\frac{\partial}{\partial \lambda} \mathcal{L} = 0 \Leftrightarrow \lambda = 1/\bar{x}$$

$$\frac{\partial^2}{\partial \lambda^2} \mathcal{L} = -\frac{n}{\lambda^2} < 0$$

$$\Rightarrow \hat{\lambda}_{MLE} = 1/\bar{x}$$

$$3) X \sim \mathcal{U}[0, \theta] \quad , \quad h(x, \theta) = \frac{1}{\theta} \mathbb{1}_{[0, \theta]}(x)$$

$$L(\theta, x_1, x_2, \dots, x_n) = \prod_{i=1}^n \frac{1}{\theta} \mathbb{1}_{[0, \theta]}(x_i)$$

$$= \frac{1}{\theta^n} \cdot \prod_{i=1}^n \mathbb{1}_{\{x_i \geq 0\}} \cdot \mathbb{1}_{\{x_i \leq \theta\}} \rightarrow \mathbb{1}_{[A \cap B]} = \mathbb{1}_A \cdot \mathbb{1}_B$$

$$= \frac{1}{\theta^n} \cdot \underbrace{\prod_{i=1}^n \mathbb{1}_{\{x_i \geq 0\}}}_{= 1 \text{ o.s.}} \cdot \prod_{i=1}^n \mathbb{1}_{\{x_i \leq \theta\}}$$

$$= \frac{1}{\theta^n} \cdot \mathbb{1}_{\{\max x_i \leq \theta\}}$$



$$\Rightarrow \hat{\theta}_{MLE} = X_{(n)} \\ = \max_i X_i$$

But Moments Methods.

- 1) $\mu_1 = E_{\theta}[X] = \frac{\theta}{2}$
- 2) $\theta = 2\mu_1$
- 3) $\hat{\theta} = 2 \cdot \bar{X}$

→ Moments Methods give the different estimator from MLE.

4) $X \sim \mathcal{N}(\mu, \sigma)$, $\theta = (\mu, \sigma)$

$$h(x, \theta) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}} \rightarrow \text{w.r.t } \mu \text{ and } \sigma$$

$$L(\theta, X_1, \dots, X_n) = \prod_{i=1}^n \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(X_i - \mu)^2}{2\sigma^2}}$$

$$= \left(\frac{1}{\sigma\sqrt{2\pi}} \right)^n \cdot e^{-\sum_{i=1}^n \frac{(X_i - \mu)^2}{2\sigma^2}}$$

$$\underbrace{LL(\theta, X_1, \dots, X_n)}_{f(\theta)} = n \cdot \ln \left(\frac{1}{\sigma\sqrt{2\pi}} \right) - \sum_{i=1}^n \frac{(X_i - \mu)^2}{2\sigma^2}$$

$$= -n \ln \sigma - n \ln(\sqrt{2\pi}) - \sum_{i=1}^n \frac{(X_i - \mu)^2}{2\sigma^2}$$

$$\nabla f: \frac{\partial f}{\partial \mu}(\theta) = \frac{1}{\sigma^2} \cdot \sum_{i=1}^n (X_i - \mu) = \frac{1}{\sigma^2} (n\bar{X} - n\mu) = \frac{n}{\sigma^2} (\bar{X} - \mu)$$

$$\begin{aligned}\frac{\partial f}{\partial \sigma}(\theta) &= -\frac{n}{\sigma} + \frac{1}{\sigma^3} \cdot \sum_{i=1}^n (x_i - \mu)^2 \\ &= \frac{n}{\sigma} \left[-1 + \frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2 \cdot \frac{1}{\sigma^2} \right]\end{aligned}$$

$$\bullet \nabla f(\theta) = 0 \Leftrightarrow \begin{cases} \mu = \bar{X} \\ -1 + \frac{1}{n} \cdot \sum_{i=1}^n (x_i - \mu)^2 \cdot \frac{1}{\sigma^2} = 0 \end{cases}$$

$$\Leftrightarrow \begin{cases} \mu = \bar{X} \\ \sigma = \sqrt{\frac{1}{n} \cdot \sum_{i=1}^n (x_i - \bar{X})^2} \end{cases}$$

$$\bullet Hf(\theta): \quad \frac{\partial^2 f}{\partial \mu \partial \sigma}(\theta) = \frac{-2}{\sigma^3} n \cdot (\bar{X} - \mu)$$

$$\frac{\partial^2 f}{\partial \mu^2}(\theta) = \frac{-n}{\sigma^2}$$

$$\frac{\partial^2 f}{\partial \sigma^2}(\theta) = \frac{n}{\sigma^2} - \frac{3}{\sigma^4} \sum_{i=1}^n (x_i - \mu)^2 \rightarrow \text{Pm các trị năng} < 0$$

Claim: $Hf(\bar{X}, \underbrace{\sqrt{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{X})^2}}_{\hat{\sigma}}) < 0 \leftarrow \text{Chỉ cần check tại } \hat{\theta}.$

If $\mu = \bar{X}, \sigma = \hat{\sigma}$

$$\begin{aligned}\frac{\partial^2}{\partial \sigma^2} f(\theta) &= \frac{n}{\hat{\sigma}^2} - \frac{3}{\hat{\sigma}^4} \underbrace{\frac{1}{n} \sum_{i=1}^n (x_i - \bar{X})^2}_{\hat{\sigma}^2} \\ &= \frac{n}{\hat{\sigma}^2} - \frac{3n}{\hat{\sigma}^2} = \frac{-2n}{\hat{\sigma}^2}\end{aligned}$$

$$\Rightarrow \text{Hff}(\bar{x}, \hat{\sigma}) = \begin{pmatrix} -\frac{n}{\hat{\sigma}^2} & 0 \\ 0 & -\frac{2n}{\hat{\sigma}^3} \end{pmatrix} : \text{negative definite}$$

Hence $\hat{\theta}_{\text{MLE}} = (\bar{x}, \hat{\sigma})$

Theorem In an identifiable and dominated statistical model, if:

1) $\theta \mapsto \ln(h(x, \theta))$ is continuous for all x .

2) $\mathcal{H}$ is a compact subset of $\mathbb{R}^k$

Then for all $\theta \in \mathcal{H}$, $\hat{\theta}_{\text{MLE}} \xrightarrow[n \rightarrow +\infty]{P_{\theta} \text{ a.s.}} \theta$.

Remarks

1) There exists a kind of CLT for MLE.

2) In practice, if $\mathcal{H}$ is not compact, it is possible to apply this Theorem to an increasing sequence of compact sets increasing to $\mathcal{H}$ and finally, get the same result.

3) Usually MLE are biased.

But under some assumptions it is possible to show that MLE are asymptotically unbiased.

4. Fisher information, Cramer Rao Lower Bound and Efficiency.

4.1) Fisher information.

Let $P_{(\Theta)} = \{P_{\theta}, \theta \in \Theta\}$ be a parametric and dominated statistical model.

Let $\hat{\lambda}$ be an estimator of $\lambda = g(\theta)$.

Def. A model is said to be regular iff

1) for almost all x , $\theta \mapsto h(x, \theta)$ is differentiable.

2) It is possible to invert $\frac{\partial}{\partial \theta}$ and $\int$ in $\frac{\partial}{\partial \theta} E_{\theta}[h(\theta, x)]$

3) $I(\theta) = E_{\theta} \left[\left(\frac{\partial}{\partial \theta} \ln h(x, \theta) \right)^2 \right]$ exists and has an inverse.

↑ Nếu θ là vector thì đây là ma trận.

$I_{ij}(\theta) = E_{\theta} \left[\frac{\partial}{\partial \theta_i} \ln h(x, \theta) \cdot \frac{\partial}{\partial \theta_j} \ln h(x, \theta) \right] \rightarrow I_{ij}(\theta)$ has an inverse
↑
hàng phi của ma trận I .

Def. When the integral exists, the Fisher Information of the model is

• for one value x

$$I(\theta) = E_{\theta} \left[\left(\frac{\partial}{\partial \theta} \ln h(x, \theta) \right)^2 \right]$$

• for a random sample $X_1, X_2, \dots, X_n$:

$$I_n(\theta) = E_{\theta} \left[\left(\frac{\partial}{\partial \theta} \ln L(\theta, X_1, \dots, X_n) \right)^2 \right]$$

Remarks :

1) $I_n(\theta) = E_{\theta} \left[\left(\frac{\partial}{\partial \theta} LL(\theta, X_1, \dots, X_n) \right)^2 \right]$

2) When θ is a vector $I(\theta)$ is a matrix

3) $I_n(\theta) = n \cdot I(\theta)$

Example $X \sim \mathcal{E}(\lambda)$

$$LL(\theta, X) = \ln(\lambda) - \lambda \cdot X$$

$$\frac{\partial}{\partial \lambda} LL(\theta, X) = \frac{1}{\lambda} - X$$

$$I(\lambda) = E_{\theta} \left[\left(\frac{\partial}{\partial \lambda} LL(\theta, X) \right)^2 \right]$$

$$= E_{\theta} \left[\left(\frac{1}{\lambda} - X \right)^2 \right]$$

$$= E_{\theta} \left[\frac{1}{\lambda^2} - 2 \cdot \frac{1}{\lambda} \cdot X + \bar{X}^2 \right]$$

$$= \frac{1}{\lambda^2} - \frac{2}{\lambda} E_{\theta}[X] + E_{\theta}[\bar{X}^2]$$

$$= \frac{1}{\lambda^2} - \frac{2}{\lambda} \cdot \frac{1}{\lambda} + \frac{2}{\lambda^2} = \frac{1}{\lambda^2}$$

$$\Rightarrow \text{Hence } I_n(\lambda) = n \cdot I(\lambda) = \frac{n}{\lambda^2}$$

• Proposition. Fisher information does not depend on dominating measure Q .

• Proposition. Under the following assumptions:

$$\left. \begin{array}{l} 1) \frac{\partial}{\partial \theta} \int h(x, \theta) dQ(x) = \int \frac{\partial}{\partial \theta} h(x, \theta) dQ(x) \\ 2) \frac{\partial^2}{\partial \theta^2} \int h(x, \theta) dQ(x) = \int \frac{\partial^2}{\partial \theta^2} h(x, \theta) dQ(x) \end{array} \right\} \text{ this can check etc.}$$

We get

$$E_{\theta} \left[\frac{\partial}{\partial \theta} \ln(h(X, \theta)) \right] = 0$$

and

$$I(\theta) = -E_{\theta} \left[\frac{\partial^2}{\partial \theta^2} \ln(h(X, \theta)) \right]$$

4.2. Cramér - Rao Blackwell Lower Bound.

Let T be a statistic such that

$$1) E_{\theta}[T^2] < +\infty$$

$$2) T \text{ is an unbiased estimator of } \lambda = g(\theta) : E_{\theta}[T] = \lambda.$$

Assume the model is regular then we get the following theorem:

• Theorem Cramér - Rao inequality.

Under previous assumptions:

$$\text{Var}_{\theta}(T) \geq \frac{|g'(\theta)|^2}{I_{\theta}(\theta)} \quad \text{if } \theta \in \mathbb{R} \\ \lambda \in \mathbb{R}.$$

$$\text{Var}_\theta(T) \geq J_g(\theta) \cdot I_n^{-1}(\theta) \cdot J_g(\theta)^T \text{ if } \theta, \lambda \text{ is vector.}$$

lay VT - VP de ma tran xet +.

where $g: \mathbb{R}^k \rightarrow \mathbb{R}^m$

$$g(\theta) = (g_1(\theta), \dots, g_n(\theta))$$

$$J_g(\theta) = \begin{pmatrix} \frac{\partial g_1(\theta)}{\partial \theta_1} & \dots & \frac{\partial g_1(\theta)}{\partial \theta_j} \\ \vdots & \ddots & \vdots \\ \frac{\partial g_n(\theta)}{\partial \theta_1} & \dots & \frac{\partial g_n(\theta)}{\partial \theta_j} \end{pmatrix}$$

Remarks

1) When $g(\theta) = \theta$ then we get $\text{Var}_\theta(T) \geq \frac{1}{I_n(\theta)}$

2) Because $I_n(\theta) = n \cdot I(\theta)$, we get:

$$\text{Var}(T) \geq \frac{|g'(\theta)|^2}{n \cdot I(\theta)} = \frac{\text{Constant}}{n}$$

4.3. Efficiency.

Def. When the assumptions of the previous theorem hold the estimator

T is said to be **efficient** iff the lower bounded is achieved:

$$\text{Var}_\theta(T) = \frac{|g'(\theta)|^2}{I_n(\theta)}$$

T is hence an **unbiased estimator** of $\lambda = g(\theta)$ with minimal variance

- **Best estimator** of λ (with respect to the quadratic risk).

Theorem Under the previous regularity assumptions the statistic T is an efficient estimator of $\lambda = g(\theta)$ if:

$$\frac{\partial}{\partial \theta} \ln L_n(\theta, X_1, \dots, X_n) \stackrel{\text{Boas}}{=} c(\theta) \cdot [T(X) - g(\theta)]$$

Back to example: 1) $X \sim \text{Ber}(\theta)$.

$$\frac{\partial}{\partial \theta} \ln L_n(\theta, X_1, \dots, X_n) = \underbrace{\frac{n}{\theta(1-\theta)}}_{c(\theta)} \underbrace{[\bar{X} - \theta]}_{\substack{T(X) \\ \downarrow \\ g(\theta)}} \rightarrow \bar{X} : \text{efficient.}$$

Corollary. If the model is regular and the MLE exists, then if T is efficient unbiased estimator of λ then it is the MLE.

5. Exhaustive statistics ^(a.i.) (sufficient statistic).

Def. A statistic T is said to be exhaustive if the conditional distribution of $(X_1, \dots, X_n)$ given T does not depend on θ .

Example. Let $(X_1, \dots, X_n) \stackrel{\text{i.i.d.}}{\sim} \mathcal{P}(\theta)$.

$$T = \sum_{i=1}^n X_i$$

We know that $T \sim \mathcal{P}(n\theta)$.

$$\begin{aligned} P_\theta((X_1, \dots, X_n) = (x_1, \dots, x_n) \mid T=t) \\ = P_\theta(\text{—————}, T=t) / P_\theta(T=t). \end{aligned}$$

$$\begin{aligned}
&= \begin{cases} \frac{P_0((x_1, \dots, x_n) = (x_1, \dots, x_n))}{P(G=t)} & \text{if } t = \sum_{i=1}^n x_i \\ 0 & \text{if } t \neq \sum_{i=1}^n x_i \end{cases} \\
&= \begin{cases} \frac{\prod_{i=1}^n \frac{\theta^{x_i}}{x_i!} \cdot e^{-\theta}}{(\sum_{i=1}^n \theta)^{\sum_{i=1}^n x_i} \cdot e^{-\sum_{i=1}^n \theta}} & \text{if } t = \sum_{i=1}^n x_i \\ 0 & \text{else} \end{cases} \\
&= \begin{cases} \frac{\cancel{\theta^{\sum_{i=1}^n x_i}}}{\prod x_i!} \cdot \cancel{e^{-\theta}} / \frac{(\sum_{i=1}^n \theta)^{\sum_{i=1}^n x_i} \cdot \cancel{e^{-\sum_{i=1}^n \theta}}}{(\sum x_i)!} & \text{if } t = \sum x_i \\ 0 & \text{else} \end{cases} \\
&= \begin{cases} \frac{(\sum x_i)!}{\prod x_i! \cdot \theta^{\sum x_i}} & \text{if } t = \sum x_i \\ 0 & \text{else} \end{cases}
\end{aligned}$$

does not depend on $\theta \rightarrow$ exhaustive.

21.10.2022.

Theorem: Let T be a statistic computed from the sample $(X_1, \dots, X_n)$.

T is said to be **exhaustive** if and only if the likelihood of the sample may be factorized in the following way.

$$L(\theta, X_1, \dots, X_n) = \psi(T, \theta) \cdot \ell(X_1, \dots, X_n).$$

Remarks

1) Another way to define an exhaustive statistic is the measurable function $E_\theta[h(X_1, \dots, X_n) | T]$ does not depend on θ .

2) If there exists an exhaustive statistic T given by the factorization theorem then the **MLE estimator** is a function of T .

Examples 1) $X_1, \dots, X_n \stackrel{i.i.d.}{\sim} \mathcal{N}(\theta, 1)$, $T = \sum_{i=1}^n X_i$

$$\begin{aligned} L(\theta, X_1, X_2, \dots, X_n) &= \prod_{i=1}^n \frac{1}{\sqrt{2\pi}} \cdot e^{-\frac{(X_i - \theta)^2}{2}} \\ &= \left(\frac{1}{\sqrt{2\pi}}\right)^n \cdot e^{-\sum_{i=1}^n (X_i^2 + \theta^2 - 2X_i\theta)/2}. \end{aligned}$$

$$\begin{aligned}
 &= (\sqrt{2\sigma})^{-n} \cdot e^{-1/2 [\sum_{i=1}^n x_i^2 + n\theta^2 - 2\theta \cdot \sum x_i]} \\
 &= \underbrace{(\sqrt{2\sigma})^{-n} \cdot e^{-\frac{1}{2} \sum_{i=1}^n x_i^2}}_{\ell(x_1, \dots, x_n)} \cdot \underbrace{e^{-\frac{1}{2} n\theta^2 + \theta \tau}}_{\psi(\tau, \theta)}.
 \end{aligned}$$

hence T is exhaustive.

2) $X_1, \dots, X_n \stackrel{i.i.d}{\sim} \mathcal{N}(\mu, \sigma)$, (θ, σ)

$$\begin{aligned}
 L(\theta, X_1, \dots, X_n) &= \prod_{i=1}^n \frac{1}{\sqrt{2\pi} \cdot \sigma} \cdot e^{-(x_i - \theta)^2 / 2\sigma^2} \\
 &= (\sqrt{2\pi} \cdot \sigma)^{-n} \cdot e^{-\sum (x_i^2 - 2x_i\theta + \theta^2) / 2\sigma^2} \\
 &= (\sqrt{2\pi})^{-n} \cdot \sigma^{-n} \cdot e^{-(\sum x_i^2 - 2\theta \sum x_i + n\theta^2) / 2\sigma^2} \\
 &= (\sqrt{2\pi})^{-n} \cdot \sigma^{-n} \cdot e^{-\sum x_i^2 / 2\sigma^2} \cdot e^{n\theta \bar{X} / \sigma^2} \cdot e^{-n\theta^2 / 2\sigma^2} \\
 &= \psi((\theta, \sigma), T),
 \end{aligned}$$

where $T = \begin{pmatrix} \bar{X} \\ \sum x_i^2 \end{pmatrix}$. We have seen that: $\hat{\theta}_{MLE} = \bar{X}$

$\hat{\sigma}_{MLE} = \sqrt{\sum x_i^2 - n(\bar{X})^2}$
 $\hookrightarrow$ function of T .

3) $X_1, \dots, X_n \stackrel{i.i.d}{\sim} \mathcal{U}[0, \theta]$, $\theta > 0$.

$$\begin{aligned}
 L(\theta, X_1, \dots, X_n) &= \prod_{i=1}^n \frac{1}{\theta} \cdot \mathbb{1}_{[0, \theta]}(x_i) \\
 &= \frac{1}{\theta^n} \cdot \underbrace{\mathbb{1}_{\{\max X_i \leq \theta\}} \cdot \mathbb{1}_{\{\min X_i \geq 0\}}}_{\psi(T, \theta)}
 \end{aligned}$$

$T = \max X_i$ is an exhaustive statistic.

Theorem Rao - Blackwell.

Let $\hat{\lambda}$ be an unbiased estimator of λ and assume $E_{\theta}[\hat{\lambda}] < +\infty$.
If there exists an exhaustive T then $\hat{\lambda}_T = E_{\theta}[\hat{\lambda}|T]$ is a better estimator of λ than $\hat{\lambda}$ for the quadratic risk

Remarks:

- $\hat{\lambda}_T$ is the projection of $\hat{\lambda}$ on the set of measurable functions of T . If $\hat{\lambda}$ is already a measurable function T then $\hat{\lambda}_T = E[\hat{\lambda}|T] = \hat{\lambda}$.
non dist. $E[\hat{\lambda}|T] = g(T)$
- $E[\hat{\lambda}_T] = E[E[\hat{\lambda}|T]] = E[\hat{\lambda}] = \lambda$ if $\hat{\lambda}$ is unbiased.

Example.

1) $X_1, \dots, X_n \stackrel{i.i.d}{\sim} \text{Ber}(\theta)$, $\theta \in]0, 1[$.

$\hat{\theta} = X_1$: unbiased estimator of θ .

$$L(\theta, X_1, \dots, X_n) = \theta^{n\bar{X}} \cdot (1-\theta)^{n-n\bar{X}}$$

Hence $T = \bar{X}$ is an exhaustive statistic.

Now the distribution for all t of $X_1|T=t$ is $\text{Ber}(t)$

Then $E[X_1|t] = T = \bar{X}$. Indeed:

$$P(X_1 = u_1 | T = t) = \frac{P(X_1 = u_1, T = t)}{P(T = t)}$$

$$= \frac{P(X_1 = x_1, \sum X_i = nt)}{P(\sum X_i = nt)}$$

$$= \sum_{x_2, \dots, x_n} \frac{P(X_1 = x_1, X_2 = x_2, \dots, X_n = x_n, \sum X_i = nt)}{P(\sum X_i = nt)}$$

$$= \sum_{\substack{x_2, \dots, x_n \\ \sum_{i=2}^n x_i = nt - x_1}} \frac{\theta^{x_1} \cdot (1-\theta)^{n-x_1} \cdot P(\sum X_i = nt)}{\binom{nt}{n} \cdot \theta^{nt} \cdot (1-\theta)^{n-nt}}$$

Bin(θ).

$$\frac{(n-1)!}{(nt-1)!(n-nt)!} \cdot \frac{n!}{(nt)!(n-nt)!} = \begin{cases} \binom{nt}{n-1} / \binom{nt}{n} & \text{if } x_1 = 0 \\ \binom{nt-1}{n-1} / \binom{nt}{n} & \text{if } x_1 = 1. \end{cases}$$

$$\frac{(n-1)!}{(nt-1)!(n-nt)!} \cdot \frac{n!}{(nt)!(n-nt)!} = \begin{cases} 1-t & \text{if } x_1 = 0 \\ t & \text{if } x_1 = 1. \end{cases}$$

$$\text{Hence } P(X_1 = x_1 | T = t) = \begin{cases} 1-t & \text{if } x_1 = 0 \\ t & \text{if } x_1 = 1. \end{cases}$$

$$\Rightarrow X_1 | T=t \sim \text{Ber}(t).$$

$$2) X_1, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} \mathcal{BP}(\theta, 1)$$

X_1 is an unbiased estimator of θ .

$$E[X_1] = \theta.$$

$\bar{x}$ exhaustive

$$\hookrightarrow E[x_i | \bar{x}] = \varphi(\bar{x}) = E[x_i | \bar{x}] \text{ (Idea)}$$

$$\varphi(\bar{x}) = \frac{1}{n} \sum_{i=1}^n E[x_i | \bar{x}] = E\left[\frac{1}{n} \sum_{i=1}^n x_i | \bar{x}\right] = E[\bar{x} | \bar{x}] = \bar{x}$$

(vi hat $\hat{\alpha}$ $E[x_i | \bar{x}] = n \cdot \bar{x}$)

$\Rightarrow E[x_i | \bar{x}] = \bar{x}$ is better estimator of θ .

6.7 Exponential families

Def. A family $\mathcal{P}_{\Theta} = \{P_{\theta}, \theta \in \Theta\}$ is an exponential family if the density of P_{θ} is the following form:

$$h(x, \theta) = e^{\langle \alpha(\theta), T_{\theta}(x) \rangle + \beta(\theta)} \cdot l(x)$$

where

$$\begin{cases} \alpha: \Theta \rightarrow \mathbb{R}^k \\ \beta: \Theta \rightarrow \mathbb{R} \\ T_{\theta}: \mathcal{X} \rightarrow \mathbb{R}^k \\ \langle \cdot, \cdot \rangle: \mathbb{R}^k \text{ inner product.} \end{cases}$$

Property: If $X_1, X_2, \dots, X_n \stackrel{i.i.d.}{\sim} P_{\theta}$, with $P_{\theta} \in \mathcal{P}_{\Theta}$ and $\mathcal{P}_{\Theta}$ is an exponential family then $T = \frac{1}{n} \sum_{i=1}^n T_{\theta}(x_i)$ is an exhaustive statistic.

Proof

$$L(\theta, x_1, \dots, x_n) = \prod_{i=1}^n e^{\langle \alpha(\theta), T_{\theta}(x_i) \rangle + \beta(\theta)} \cdot l(x_i)$$

$$\begin{aligned}
&= \left[\prod_{i=1}^n l(x_i) \right] \cdot e^{\sum_{i=1}^n [\langle \alpha(\theta), T_0(x_i) \rangle + \beta(\theta)]} \\
&= \prod_{i=1}^n l(x_i) \cdot e^{\langle \alpha(\theta), \sum_{i=1}^n T_0(x_i) \rangle + n\beta(\theta)} \\
&= \prod_{i=1}^n l(x_i) \cdot e^{\underbrace{\langle \alpha(\theta), nT \rangle + n\beta(\theta)}_{\psi(T, \theta)}}
\end{aligned}$$

↳ T is exhaustive. $\square$

Theorem. Lehman Scheffe.

If the model is exponential and $\alpha(\theta)$ contains an open subset of $\mathbb{R}^k$, then for any function $f: \mathbb{R}^k \rightarrow \mathbb{R}^1$ such that $\lambda = E_\theta[\|f(T)\|] < +\infty$, the statistic $f(T)$ is the best unbiased estimator of $E_\theta[f(T)]$.

Comment

If you want to estimate $\lambda = g(\theta)$ under the assumption of the previous theorem, if you define an estimator $\hat{\lambda}$ of λ that is unbiased and is written as a function of T (with $E[\|\hat{\lambda}\|] < +\infty$) then $\hat{\lambda}$ is the best estimator of λ .

Def. A statistic T is said to be complete iff.

$$\forall \theta \in \mathcal{H}, \forall g \text{ borel measurable function} \\
(E_\theta[g(T)] = 0) \Rightarrow P_\theta(g(T) = 0) = 1.$$

Property. In an exponential family, if $\alpha(\cdot)$ contains a non empty open subset of $\mathbb{R}^k$, then T is complete.
 $\downarrow$ $\uparrow$ exhaustive & free.

Idea:

$$\begin{aligned} E_0[g(T)] &= \int g(T) \cdot e^{\langle \alpha(\theta), T \rangle + \beta(\theta)} \cdot l(x_1, \dots, x_n) dQ(x_1) dQ(x_2) \dots dQ(x_n) \\ &= \int g(t) \cdot e^{\langle \alpha(\theta), t \rangle + \beta(\theta)} \cdot \frac{l(x_1, \dots, x_n) \cdot dQ(x_1, \dots, x_n | T=t)}{dQ_0(t)} \\ &= \int \underbrace{g(t) \cdot p(t)}_{m(t)} \cdot e^{\langle \alpha(\theta), t \rangle + \beta(\theta)} dQ_0(t) \cdot p(t) \\ &= \int m(t) \cdot e^{\langle \alpha(\theta), t \rangle} dQ_0(t). \end{aligned}$$

Property of Laplace transform, if Λ has non empty open subset that:

$$\left[\text{for all } \lambda \in \Lambda, \int e^{\langle \lambda, t \rangle} m(t) dQ_0(t) = 0 \right] \Rightarrow m=0 \quad Q_0 \text{ a.s.}$$

Theorem. Lehman - Scheffe.

If T is an exhaustive and complete statistic

$\hat{\lambda}$ is an unbiased estimator of λ .

then $\hat{\lambda}_T = E[\hat{\lambda} | T]$ is a uniform minimum variance unbiased estimator of λ .