

Statistical inference
II

Estimation
Hypothesis

Book

1. Introduction to theory of statistics - Mood and Graybill
2. Statistical Inference $\rightarrow$ Casella and Berger
3. Introduction to Mathematical statistics -
Hogg and Craig

Modal and Median unbiased estimator:

$$\text{Statistic} = 5.5 \rightarrow \text{Estimator} = \frac{\sum_{i=1}^n x_i}{n}$$

$\Rightarrow$ estimate $5.5 + 1 = 6.5$

Qn. difference among estimator, statistic & estimate

H.W. Die $\rightarrow$ toss, find the median or modal unbiased estimator.
Outcomes = 1, 2, 3, 4, 5, 6

$X \rightarrow$ random variable

Hint: $X \Rightarrow x_1, x_2, x_3, x_4, x_5, x_6$
 $P(X = x_1) = 1/6, \dots, P(X = x_6) = 1/6 \rightarrow$ follows discrete uniform distⁿ

According to question, the distⁿ of random variable is defined as discrete uniform distribution.

$$P(X; \theta) = \frac{1}{\theta} ; \quad x = 0, \dots, \theta$$

Here, $\theta = 6$

~~Note~~ Mean = $E(X)$

$$= \sum x P(x)$$

$$= 1 \cdot P(X=1) + 2 \cdot P(X=2) + \dots + 6 \cdot P(X=6)$$

$$= 1 \cdot \frac{1}{6} + 2 \times \frac{1}{6} + 3 \times \frac{1}{6} + 4 \times \frac{1}{6} + 5 \times \frac{1}{6} + 6 \times \frac{1}{6}$$

$$= \frac{1}{6} + \frac{2}{3} + \frac{1}{2} + \frac{2}{3} + \frac{5}{6} + 1$$

$$= \frac{1+2+3+4+5+6}{6}$$

$$= \frac{21}{6}$$

H.W. $\frac{7}{2}$ Die $\rightarrow$ toss, giving the 5-3 regions on number line

Outcomes = 1, 2, 3, 4, 5, 6

Random variable $\rightarrow X$

Hint: $X \Rightarrow$... $P(X=x) = \frac{1}{6}$... $P(X=6) = \frac{1}{6}$

* Let x is a r.v for the value $1, 2, \dots, 6$

$$P(X = x_1) = 1/6, \quad P(X = x_6) = 1/6$$

which follows discrete distⁿ

$$\text{Mean, } E(X) = \sum_{x=1}^6 x \cdot P(x) = 1 \times \frac{1}{6} + \frac{2}{6} + \frac{3}{6} + \frac{4}{6} + \frac{5}{6} + 1$$

$$= \frac{21}{6}$$

$$= 3.5$$

$\therefore$ Sample mean is an unbiased estimator of popⁿ mean, $E(\bar{X}) = \mu = 3.5$

Median: Since the outcomes are 1, 2, 3,

4, 5, 6. The median is the average

of 3rd and 4th numbers

$$\text{median} = \frac{3+4}{2} = 3.5$$

Mode: As the mode is the most frequent value and for a fair die

to toss all outcomes are equally likely. $\therefore$ (uniform distⁿ) $\therefore$ So there is no unique mode.

Median or mode unbiased estimator

Die $\rightarrow \{1, 2, 3, 4, 5, 6\} \rightarrow$ discrete

median = ?

$E[X] = ?$

If $E(X) = \text{median} \rightarrow$ median unbiased

If $E(X) = \text{mode} \rightarrow$ mode unbiased

Estimator = $\bar{x} = \frac{\sum x_i}{n}$

Statistic

$$\mu = \bar{X} = \frac{X_1 + X_2 + \dots + X_n}{N}$$

Estimator but statistic is
করা অবশ্যই popⁿ

If sample নিয়ে কাজ করি তাহলে

আমি estimator and statistic

শেষ

Properties of good estimator:

1. Unbiasedness $\rightarrow E[\text{Statistic}] = \text{Pop}^n \text{ Parameter}$
 $\downarrow$
 estimator $\rightarrow$ unbiased estimator
2. Consistency
3. Sufficiency
4. Efficiency

$$E[\bar{x}] = \mu$$

$$E(s^2) = \sigma^2$$

HW $s^2 = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n-1}$ is a biased estimator of popⁿ variance σ^2

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

$$= \frac{1}{n} \sum (x_i - \mu)^2 - \frac{2}{n} \sum (x_i - \mu)(\bar{x} - \mu) + \frac{1}{n} \sum (\bar{x} - \mu)^2$$

$$= \frac{1}{n} \sum (x_i - \mu)^2 - 2(\bar{x} - \mu) \sum \frac{1}{n} + (\bar{x} - \mu)^2$$

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

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

$$= \frac{1}{n} \sum \text{var}(x_i) - \text{var}(\bar{x})$$

$$= \frac{1}{n} \sum_{i=1}^n \sigma^2 - \frac{\sigma^2}{n}$$

$$= \left(\frac{n-1}{n} \right) \sigma^2$$

$\therefore S^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$ is ~~a~~ biased estimator of σ^2

$$E(S^2) = \left(\frac{n-1}{n} \right) \sigma^2$$

$$\Rightarrow E\left[\frac{n}{n-1} \cdot \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2\right] = \sigma^2$$

$$E\left[\frac{n}{n-1} \cdot \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2\right] = \sigma^2$$

$$\Rightarrow E\left[\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2\right] = \sigma^2$$

So, $\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$ is an unbiased estimate of σ^2



Amount of bias = Estimate - Parameter

$$= \left(\frac{n-1}{n}\right) \sigma^2 - \sigma^2$$

$\therefore -\frac{\sigma^2}{n}$ is the amount of the bias

Lecture 5th & 6th

22-06-25

⊗ Properties of good estimator

(i) Unbiasedness

(ii) Consistency

(iii) Efficiency

(iv) Sufficiency

Ex: If X_i is a Bernoulli random variable with parameter p , then:

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

is the maximum likelihood estimator (MLE) of p . Is the MLE of p an unbiased estimator of p ?

$$f(x, \lambda) = \log \left[\frac{e^{-\lambda} \lambda^x}{x!} \right]$$

$$\frac{\partial}{\partial \lambda} \left[\log \left(\frac{e^{-\lambda} \lambda^x}{x!} \right) \right] = \exp \left[\frac{\partial \left[\log \left(\frac{e^{-\lambda} \lambda^x}{x!} \right) \right]}{\partial \lambda} \right]$$

$$\text{var}(\hat{\lambda}) = \frac{1}{n} \text{var} \left[\sum_{i=1}^n X_i \right]$$

$$= \frac{n \lambda}{n} = \frac{\lambda}{n}$$

same as

Binomial

$$P[(0, 1, 1) | S] = ?$$

$$P[(0, 1, 1) | T] = ?$$

$$P = [X_1 = 0, X_2 = 1, X_3 = 1 | S = 2]$$

$$= \frac{P^1(1-P)}{P(S=2)}$$

$$= \frac{P^1(1-P)}{\binom{3}{2} \cdot P^2(1-P)} = \frac{1}{3}$$

$$P(0, 1, 1 \mid 1) = \frac{P[X_1=0, X_2=1, X_3=1; T=1]}{P[T=1]}$$

$$= \frac{P(1-p)^2 p}{(1-p)^2 p} = 1$$

$$\frac{P}{1+P}$$

$$\frac{1}{2} = (0; x)$$

* Median & Unbiased Estimator

Let, $X_1, X_2, \dots, X_n$ be iid random variables with pdf $f(x; \theta)$ and let $t(X_1, \dots, X_n)$ be a statistic such that the median of the pdf of the statistic t is θ . Then the statistic $t(X_1, \dots, X_n)$ is defined to be an median estimate of θ .

$$0 = (t) - \theta \Rightarrow 0 = \text{Bias}$$

For discrete distⁿ median will be

$$\text{Median} = \sum_{x=0}^{\infty} P_n(X=x) = \frac{1}{\theta}$$

and for continuous distⁿ

$$\text{Median} = \int_{\text{min Range}}^{\text{Median}} f(x; \theta) = \frac{1}{2}$$

$$= \int_{\text{min Range}}^{\text{Max Range}} f(x; \theta) = \frac{1}{2}$$

To estimate the popⁿ parameter θ by statistic t , we need to know that about bias.

$$\text{So Bias} = \theta - E(t)$$

To make the statistic t unbiased, bias should be zero.

$$\therefore \text{Bias} = 0 \Rightarrow \theta - E(t) = 0$$

$$\Rightarrow E(t) = 0$$

$$P_n \{ E(t) > 0 \} = P_n \{ E(t) < 0 \}$$

So, if these two inequalities holds for both sides, then the statistic t will be median unbiased estimator of θ

parameter θ .

Lecture 7th & 8th

30-06-25



Modal Unbiased Estimator:

Let, $X_1, X_2, \dots, X_n$ be iid distributed random variables with probability density function $f(x; \theta)$ and let $t(X_1, X_2, \dots, X_n)$ be a statistic. Such that the mode of the density function of statistic, t is θ . Then the statistic t is said to be a modal unbiased estimate of θ .

$$E(\bar{x}) = \theta$$

↳ mode

$E(\bar{x}) = \theta$ median

H-W

Example:

Let $X_1, \dots, X_n$ be iid random samples from a continuous uniform distⁿ

$$f(x; \theta) = \begin{cases} \frac{1}{\theta} & 0 \leq x \leq \theta \\ 0 & \text{otherwise} \end{cases}$$

Find the modal unbiased estimator of the popⁿ parameter θ .

→ Example

Ex: Let $X_1, \dots, X_n$ be iid random samples from an continuous exponential distⁿ.

$$f(x; \theta) = \frac{1}{\theta} e^{-x/\theta} ; \quad x \geq 0, \theta > 0$$

Let the two statistic,

$$Y_1 = \frac{X_n}{\log n} \quad \text{where} \quad X_n = \max \{X_1, X_2, \dots, X_n\}$$

$$\text{and } Y_2 = \frac{\delta_n}{n-1} \quad \text{where} \quad \delta_n = \sum_{i=1}^n X_i$$

Prove that both statistics Y_1 and Y_2 are the modal unbiased estimate of the

popⁿ parameter θ .

Solⁿ: The prob^y density function of the highest order statistic, x_n is

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

$$= n \left[\int_0^{x_n} \frac{1}{\theta} e^{-x/\theta} dx \right]^{n-1} \cdot \frac{1}{\theta} e^{-x_n/\theta} ; x_n > 0$$

$$= n \cdot \frac{1}{\theta} \left[e^{-x/\theta} \right]_0^{x_n} \cdot \frac{1}{\theta} e^{-x_n/\theta}$$

$$= \frac{n}{\theta} \left[\frac{1}{\theta} e^{-x_n/\theta} \right]$$

$$= \frac{n}{\theta} \left[\left(\frac{1}{\theta} \cdot e^{-x_n/\theta} \right) \cdot \left(-\frac{1}{\theta} \right) \right]^{n-1} e^{-x_n/\theta}$$

$$= \frac{n}{\theta} \left[-e^{-x_n/\theta} + e^{-x_n/\theta} \right]^{n-1} \cdot e^{-x_n/\theta}$$

$$= n \frac{n}{\theta} \left[1 \right] e^{-x_n/\theta} \cdot e^{-x_n/\theta}$$

Given that,

$$y_1 = \frac{x_n}{\log n}$$

$$\Rightarrow x_n = y_1 \log n$$

$$\Rightarrow \frac{\partial x_n}{\partial y_1} = \log n$$

$$\Rightarrow |J| = \left| \frac{\partial x_n}{\partial y_1} \right| = \log n$$

$$\therefore g(y_1; \theta) = g(x_n, \theta) \cdot |J|$$

$$= \frac{n}{\theta} \left[1 - e^{-x_n/\theta} \right]^{n-1} e^{-x_n/\theta} \cdot \log n$$

$$= \frac{n \log n}{\theta} \left[1 - e^{-\frac{x_n}{\theta}} \right]^{n-1} e^{-\frac{x_n}{\theta}}$$

$$= \frac{n \log n}{\theta} \left[1 - e^{-\frac{y_1 \log n}{\theta}} \right]^{n-1} e^{-\frac{y_1 \log n}{\theta}}$$

$$= \frac{n \log n}{\theta} \left[1 - e^{-\frac{y_1 \log n}{\theta}} \right]^{n-1} e^{-\frac{y_1 \log n}{\theta}}$$

$$= \frac{n \log n}{\theta} \left[1 - e^{-y_1/\theta} \right]^{n-1} e^{-y_1/\theta}$$

$$= \frac{n \log n}{\theta} \left[1 - \frac{ny_1}{\theta} \right]^{n-1} \cdot \left(-\frac{ny_1}{\theta} \right) \Rightarrow$$

Now,

Let,

$$\frac{n-1}{n} = m$$

$$\frac{n \log n}{\theta} \left[1 - m^{y_1} \right]^{n-1} \cdot m^{y_1}$$

Differentiate both sides w.r. to y_1 and put equal to zero

$$\frac{\partial g(y_1; \theta)}{\partial y_1} = 0$$

$$\Rightarrow \frac{\partial}{\partial y_1} \left[\frac{n \log n}{\theta} \left\{ 1 - m^{y_1} \right\}^{n-1} \cdot m^{y_1} \right] = 0$$

$$\Rightarrow \frac{\partial}{\partial y_1} \left[\left(1 - m^{y_1} \right)^{n-1} \cdot m^{y_1} \right] = 0$$

$$\Rightarrow (1 - m^{y_1})^{n-1} m^{y_1} \log n + m^{y_1} (n-1) (1 - m^{y_1})^{n-2} \times$$

$$[0 - m^{y_1} \log m] = 0$$

$$\left[\frac{2}{2x} a^x = a^x \log a \right]$$

$$\Rightarrow (1 - m^{y_1})^{n-1} m^{y_1} \log m - m^{2y_1} (n-1) (1 - m^{y_1})^{n-2}$$

$$\Rightarrow \log m \left[(1 - m^{y_1})^{n-1} m^{y_1} - m^{2y_1} (n-1) (1 - m^{y_1})^{n-2} \right] = 0$$

$$\Rightarrow (1 - m^{y_1})^{n-1} m^{y_1} = m^{2y_1} (n-1) (1 - m^{y_1})^{n-2}$$

$$\Rightarrow \frac{(1 - m^{y_1})^{n-1}}{(1 - m^{y_1})^{n-2}} = \frac{(n-1) m^{2y_1}}{m^{y_1}}$$

$$\Rightarrow (1 - m^{y_1})^{n-1-n+2} = (n-1) m^{2y_1-y_1}$$

$$\Rightarrow (1 - m^{y_1}) = (n-1) m^{y_1}$$

$$\Rightarrow 1 - m^{y_1} = nm^{y_1} - m^{y_1}$$

$$\Rightarrow nm^{y_1} = 1$$

$$\Rightarrow \underline{y_1 \log nm = 0}$$

$$\Rightarrow m^{y_1} = \frac{1}{n}$$

$$\Rightarrow (n^{-1/\theta})^{y_1} = \frac{1}{n}$$

$$\Rightarrow n^{-y_1/\theta} = n^{-1}$$

$$\Rightarrow \left(1 - \frac{y_1}{\theta}\right) = 1 - \frac{1}{\theta}$$

$$\Rightarrow \hat{y}_1 = \theta$$

For y_2 ,

We know that the prob^y density

$$f(x; \alpha, \beta) = \frac{\beta^\alpha \cdot e^{-\beta x} \cdot x^{\alpha-1}}{\Gamma(\alpha)}$$

$$\therefore f(S_n; n, \theta) = \frac{1}{\Gamma(n)} \left(\frac{1}{\theta}\right)^n \cdot e^{-S_n/\theta} \cdot S_n^{n-1}$$

Given that

$$y_2 = \frac{S_n}{n-1}$$

$$\frac{\partial S_n}{\partial y_2} = n-1$$

$$\Rightarrow S_n = y_2(n-1)$$

$$\Rightarrow |J| = \left| \frac{\partial S_n}{\partial y_2} \right| = (n-1)$$

Prob^y of y_2 is

$$g(y_2; \theta) = g(S_n; n, \theta) |J|$$

$$f(S_n; n, \theta) = \frac{1}{\Gamma(n)} \left(\frac{1}{\theta}\right)^n \cdot e^{-\frac{(n-1)y_2}{\theta}} \left[y_2 (n-1)\right]^{n-1}$$

$$= \frac{1}{\Gamma(n)} \left(\frac{1}{\theta}\right)^n \cdot e^{-\frac{(n-1)y_2}{\theta}} \left[y_2 (n-1)\right]^{n-1}$$

$$= \frac{(n-1)^{n-1+1}}{\theta^n \Gamma(n)} e^{-\frac{(n-1)y_2}{\theta}} y_2^{n-1}$$

$$\frac{\partial}{\partial y_2} \{g(y_2; \theta)\} = 0$$

$$\Rightarrow \frac{\partial}{\partial y_2} \left[e^{-\frac{(n-1)y_2}{\theta}} \cdot y_2^{n-1} \right] = 0$$

$$\therefore \hat{y}_2 = \theta$$

$$(1-\alpha) \frac{z_{\alpha/2}}{\sqrt{n}} = \frac{z_{\alpha/2} \sigma}{\sqrt{n}}$$

$$1-\alpha = \frac{z_{\alpha/2}^2 \sigma^2}{n \cdot \text{L.S.G}}$$

H.W → download book - Stuart and Kendall and Ord. Keith

⊗ Resampling techniques ⊗

$$\bar{x}_n = n\bar{x}_n + (n-1)\bar{x}_{n-1}$$

$$\bar{x} = \frac{1+2+3+4}{4} = 2.5$$

$$\hat{x}_1 = \hat{\theta}_1 = 4 \times 2.5 + (4-1) \times 2$$

$$J_1 = \frac{1+2+3}{3} = 2$$

$$\hat{\theta}_2 = 4 \times 2.5 + 3 \times 2.33$$

$$J_2 = \frac{1+2+4}{3} = 2.33$$

$$\hat{\theta}_3 = 4 \times 2.5 + 3 \times 2.67$$

$$J_3 = \frac{1+3+5+1}{4} = 2.67$$

$$\hat{\theta}_4 = 4 \times 2.5 + 3 \times 3$$

$$J_4 = 3$$

950 2010 - 910 - 01 01 06 00 00 00

$$J_5 = \frac{2+5+1}{3} = 2.67$$

$$J_6 = \frac{1+3+5}{3} = 3$$

High 4-9th & 10th

10-07-25

Resampling Technique:

a) Jackknife

Ex. We have 4 original samples

R.S. = 1, 2, 3, 4

$$\bar{x} = \frac{1+2+3+4}{4} = 2.5 = \frac{\sum X_i}{N}$$

~~Jackknife estimate = $n\bar{x} - (n-1)\bar{x}_{n-1}$~~

the four delete-one values are

$$\hat{x}_4 = \frac{1+2+3}{3} = 2$$

$$\hat{x}_3 = \frac{2+3+4}{3} = 3$$

$$\hat{x}_2 = \frac{3+4+1}{3} = 2.67$$

$$\hat{x}_1 = \frac{4+1+2}{3} = 2.33$$

the following pseudo values/Jackknife samples are

$$\hat{\theta}_1 = n\bar{x} - (n-1)\bar{x}_1 = 4 \times 2.5 - 3 \times 2 = 4$$

similarly, $\hat{\theta}_2 = 1$, $\hat{\theta}_3 = 2$, $\hat{\theta}_4 = 3$ absolute value

In general, Jackknife estimate $= n\bar{x} - (n-1)\bar{x}_{n-k}$

$$\text{Mean} = \frac{\sum \hat{\theta}_i}{n} = \hat{\theta}_a$$

$$\text{var}(\hat{\theta}_a) = \frac{\text{var}(\hat{\theta}_i)}{n} = \frac{\sum (\hat{\theta}_i - \hat{\theta}_a)^2}{n(n-1)}$$

$$\text{Sd}(\hat{\theta}_a) = \sqrt{\text{var}(\hat{\theta}_a)}$$

$$\text{C.I.} = \hat{\theta}_a \pm t_{\alpha/2, n-1} \sqrt{\text{var}(\hat{\theta}_a)}$$

H.W
Ex: 3.2

17.23, 13.93, 15.78, 4.91, 18.21, 14.28, 18.83,
13.45, 18.71, 18.81, 11.29, 13.39, 11.57, 10.99,
15.52, 15.25,

$$s^2 = 7.171, \quad \bar{x} = 1.9701$$

Find out the Jackknife
samples with estimated mean
variance and 95% C.I

Q. 1

What is Jackknife estimation method?

Let, $S^2 = \frac{\sum (x_i - \bar{x})^2}{N}$ is a consistent estimator

if $n \rightarrow \infty$. Like S^2 , if we have a consistent estimator but it is biased

estimator t_n of θ then we can add a simple adjustment factor to make an unbiased estimator.

Maurice - Quenouille (1956) gives a method of overcoming this type of biasness.

This method is called Jackknife estimation method.

For example,

$$E(S^2) = \frac{1}{n} \mu_2 - \frac{1}{n^2} \mu_2$$

$$\mu_2 = \frac{1}{n} \mu_2$$

$$= \frac{1}{n} \mu_2 + \frac{1}{n} \mu_2$$

Let, t_n be a biased estimator of θ based on $X_1, X_2, \dots, X_n$ observations.

$$E(t_n) = \theta + \sum_{n=1}^{\infty} \frac{a_n}{n^n}$$

Let t_{n-1} be another biased estimate of θ based on $(n-1)$ observations.

So, according to jackknife method of

estimation the Jackknife estimate is

$$t'_n = nt_n - (n-1)t_{n-1}$$

$$= t_n - t_n + nt_n - (n-1)t_{n-1}$$

$$= t_n + t_n(n-1) - (n-1)t_{n-1}$$

$$= t_n + (n-1)[t_n - t_{n-1}]$$

$$E[t'_n] = E[t_n] + (n-1)E[t_n - t_{n-1}]$$

$$= \theta + \sum_{n=1}^{\infty} \frac{a_n}{n^n} + (n-1) \left[\theta + \sum_{n=1}^{\infty} \frac{a_n}{n^n} - \left(\theta + \sum_{n=1}^{\infty} \frac{a_n}{(n-1)^n} \right) \right]$$

$$E \left[\theta - \sum_{n=1}^{\infty} \frac{a_n}{(n-1)^n} \right]$$

$$\Rightarrow F(t_n') = \theta + \sum_{n=1}^{\infty} \frac{a_n}{n^n} + n \cdot \sum_{n=1}^{\infty} \frac{a_n}{n^n} - \sum_{n=1}^{\infty} \frac{a_n}{n^n}$$

$$(n-1) \sum_{n=1}^{\infty} \frac{a_n}{(n-1)^n}$$

$$= \theta + n \sum_{n=1}^{\infty} \frac{a_n}{n^n} - (n-1) \sum_{n=1}^{\infty} \frac{a_n}{(n-1)^n}$$

$$= \theta + \sum_{n=1}^{\infty} \frac{a_n}{n^{n-1}} - \sum_{n=1}^{\infty} \frac{a_n}{(n-1)^{n-1}}$$

$$= \theta + 0 + \left\{ \frac{a_2}{n} - \frac{a_2}{(n-1)} \right\} + \left\{ \frac{a_3}{n^2} - \frac{a_3}{(n-2)^2} \right\}$$

$$+ \dots + \left\{ \frac{a_n}{n^{n-1}} - \frac{a_n}{(n-1)^{n-1}} \right\}$$

$$\therefore E(t_n') = \theta + \frac{a_2}{n(n-1)} + \text{higher orders}$$

$$E(t_n) = \theta + \sum_{n=1}^{\infty} \frac{a_n}{n^n}$$

$$E(t_n') = n t_n' - (n-1) t_{n-1}'$$

H.W
200 200
200 200

$$\left[\frac{a_n}{(n-1)^{n-1}} - \frac{a_n}{n^n} \right]$$

Example:

Find out the unbiased estimator of θ^2 in Binomial distⁿ according to the Jackknife method of estimation.

Solⁿ:

The prob^y density f^n of Binomial is

$$P_n\{x; \theta, n\} = {}^n C_x \cdot \theta^x (1-\theta)^{n-x}$$

where, $\theta = \frac{p}{n}$, $1-\theta = 1 - \frac{p}{n} = \frac{n-p}{n}$

Let, $\frac{p}{n} = \theta$ is an unbiased estimator but

$\left(\frac{p}{n}\right)^2 = \theta^2$ is not an unbiased

estimator.

Let $\left(\frac{p}{n}\right)^2$ the intuitive estimator is

$$t_n = \theta^2 = \left(\frac{p}{n}\right)^2$$

$$E[t_n] = E\left[\left(\frac{p}{n}\right)^2\right]$$

$$= \text{var}\left(\frac{p}{n}\right) + \left\{E\left(\frac{p}{n}\right)\right\}^2$$

$$V(x) = E(x^2) - \{E(x)\}^2$$

$$E(x^2) = V(x) + \{E(x)\}^2$$

$$= \frac{1}{n^2} \text{var}(n) + \theta^2$$

$$= \frac{1}{n^2} n \theta (1-\theta) + \theta^2$$

$$= \frac{\theta(1-\theta)}{n} + \theta^2$$

$$\neq \theta^2$$

According to Jackknife method the Jackknife estimate is

$$t_n' = nt_n - (n-1)\bar{t}_{n-1} \quad \text{--- (1)}$$

Now, $\bar{t}_{n-1}$ can be taken as $\left(\frac{n-1}{n-1}\right)^2$ if a

success is omitted from the sample

with prob^y $\theta = \frac{n}{n}$.

Or, $\bar{t}_{n-1}$ can be taken as $\left(\frac{n}{n-1}\right)^2$ if

a failure is omitted from the

sample with prob $= 1 - \theta = 1 - \frac{n}{n} = \frac{n-n}{n}$

$$F_{n-1} = s + f$$

$$= \left(\frac{n-1}{n}\right)^2 \cdot \frac{n}{n} + \left(\frac{n}{n-1}\right) \cdot \frac{n-n}{n}$$

$$= \frac{1}{n(n-1)^2} \left[n + n^2(n-2) \right]$$

① $\Rightarrow$

$$t'_n = nt_n - (n-1)\bar{t}_{n-1}$$

$$= n \left(\frac{n}{n}\right)^2 - (n-1) \left[\frac{1}{n(n-1)^2} \{ n + n^2(n-2) \} \right]$$

$$= \frac{n^2}{n} - \frac{1}{n(n-1)} \{ n + n^2(n-2) \}$$

$$= \frac{n(n-1)}{n(n-1)}$$

$$E[t'_n] = E \left[\frac{n(n-1)}{n(n-1)} \right]$$

Hw

Q

Given,
Number of observations:

$$n = 16$$

$$\sum x_i = \frac{17.23 + 13.93 + \dots + 10.99}{16} = 242.20$$

$$\text{Sample mean, } \bar{x} = \frac{242.20}{16} = 15.13$$

Jackknife sample means, $\bar{x}_{(i)}$

$$\bar{x}_{(i)} = \frac{\sum x_i}{15}$$

We remove each observation one by one, then calculate the sample mean of the remaining 15 observations;

Formula:

$$\bar{x}_{(i)} = \frac{\sum x - x_i}{15}$$

$\frac{i}{1}$

$\frac{x_i}{17.23}$

$$\bar{x}_{(i)} = \frac{242.20 - x_i}{15}$$
$$\frac{224.97}{15} = 14.93$$

2

13.93

$$\frac{228.27}{15} = 15.05$$

3

15.78

$$\frac{226.42}{15} = 15.08$$

i	x_i	$x_i^{(j)} = \frac{242.20 - x_i}{15}$	i
4	14.91	15.18	1
5	18.21	14.79	2
6	14.28	15.19	3
7	18.83	14.54	4
8	13.45	15.23	5
9	18.71	14.57	6
10	18.81	14.57	7
11	11.29	15.39	8
12	13.39	15.20	9
13	11.57	15.42	10
14	10.94	15.44	11
15	15.52	15.11	12
16	15.25 15.25	15.23	13

Jackknife variance

$$\text{Var}_{\text{jack}} = \frac{n-1}{n} \sum_{i=1}^n (\bar{x}_{(i)} - \bar{x})^2$$

where $\bar{x} = 15.13$

Now,

1

1

2

3

4

5

6

7

8

9

10

11

12

13

14

15

16

$\bar{x}_0$

14.33

15.05

15.08

15.18

14.79

15.19

14.54

15.25

14.57

14.57

15.39

15.20

15.42

15.44

15.11

15.13

$(x_i - \bar{x})^2$

$(14.33 - 15.13)^2 = 0.6468$

0.0068

0.0032

0.0023

0.1153

0.0038

0.3505

0.0126

0.3110

0.3198

0.0663

0.0048

0.0807

0.0952

0.0006

0.0000

$\sum_{i=1}^n \frac{1}{n}$

12.13

$$\sum (\bar{x}_i - \bar{x})^2 = 2.0199$$

$$\text{Var}_{\text{jack}} = \frac{15}{16} \times 2.0199 = 1.8937$$

$$\therefore SE = \sqrt{1.8937} = 1.3769$$

$$t_{0.025, 15} = 2.131$$

95% CI is given by

$$\begin{aligned} \text{CI} &= \bar{x} \pm t \times SE = 15.1375 \pm 2.131 \times 1.3769 \\ &= 15.1375 \pm 2.9345 \end{aligned}$$

$$= (12.20, 18.07)$$

13-07-25

Bootstrap method:

5, 9, 7, 3, 11 $\rightarrow$ resample $\rightarrow$ $\bar{x}$ $\rightarrow$ 7.8

5 times,

5 times,
 $S = n_1 =$ Bootstrap sample $o_1 = 9, 7, 9, 9, 11 = B_1 = \bar{x}_1^* = \bar{x}_1^*$

$$S = n_2 = \dots = 20 = \dots$$

$s = n_3 =$

$$S = nq =$$

$s = n_4 =$ " 2480-2 " 05 = 21 χ^2
 $s = n_5 =$

Bootstrap mean, $= \frac{1}{B} \sum_{i=1}^B \bar{X}_i^*$

" S.d =

11 S. Error: $\sqrt{\frac{\sum_{j=1}^B (\bar{X}_j^* - \bar{X})^2}{B-1}}$

11 C. T. 10

$$\rightarrow = \hat{\mu}_B \pm \frac{1}{2} \alpha_{1/2} \sqrt{\text{var}(\hat{\mu}_B)}$$

Ex → 4.1

Q consider the following dataset:

81, 70, 63, 54, 74, 65, 98, 71

- Generate 10 Bootstrap sample
- Calculate the Bootstrap Mean, S.E and 95% C.I using Mean.

Solution:

- We will randomly select 8 numbers with replacement from the dataset 10 times.

The bootstrap samples and their means are:

Sample No.	Bootstrap Sample	Mean ($\bar{x}_i^*$)
1	70, 74, 81, 54, 98, 70, 65, 74	73.25
2	54, 54, 81, 65, 63, 70, 71, 74	66.50
3	98, 98, 63, 65, 74, 71, 70, 70	76.13
4	63, 63, 65, 65, 70, 71, 74, 70	67.63
5	98, 81, 98, 74, 81, 98, 63, 54	80.88
6	54, 54, 54, 54, 81, 70, 63, 74	63.00
7	81, 81, 70, 65, 74, 71, 81, 70	74.13
8	98, 98, 98, 98, 81, 81, 81, 81	89.50
9	54, 54, 54, 54, 54, 54, 54, 54	54
10	63, 63, 63, 63, 63, 63, 63, 63	63

(ii) Bootstrap Mean

$$\begin{aligned}\bar{x}_{boot} &= \frac{1}{10} \sum_{i=1}^{10} (\bar{x}_i^*) \\ &= \frac{1}{10} (73.25 + 66.5 + 76.13 + 67.63 + 80.88 \\ &\quad + 63 + 74.13 + 89.5 + 54 + 63) \\ &= \frac{708.02}{10}\end{aligned}$$

standard error,

$$S.E = \sqrt{\frac{\sum_{i=1}^{10} (\bar{x}_i^* - \bar{x}_{boot})^2}{10-1}}$$

$$= \sqrt{\frac{(73.25 - 70.8)^2 + (66.5 - 70.8)^2 + \dots + (63 - 70.8)^2}{10-1}}$$

$$= \sqrt{\frac{927.56}{9}}$$

$$\approx 10.15$$

95% Confidence interval

$$CI = \bar{x}_{boot} \pm 1.96 \times S.E$$

$$= 70.80 \pm 1.96 \times 10.15$$

$$= 70.80 \pm 19.89$$

$$\therefore CI = [50.91, 90.69]$$

$\therefore$ Bootstrap mean = 70.80

Bootstrap standard error = 10.15

95% C.I. = [50.91, 90.69]

Bootstrap sample:

Suppose we have observed n independent random samples, say, $X = x_1, x_2, x_3, \dots, x_n$

We compute a statistic, $S(x)$

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

A bootstrap sample $x^* = x_1^*, x_2^*, \dots, x_B^*$ is obtained by randomly selecting n times with replacement from the original data, $x_1, x_2, \dots, x_n$, where each of sample size is n .

The Bootstrap estimate of the standard error is — the standard deviation of the Bootstrap samples.

$$S.E. = \sqrt{\frac{\sum_{i=1}^B [S(x_i^*) - S(x)]^2}{(B-1)}}$$

where, $S(x^*) = \frac{\sum_{i=1}^B (S(x_i^*) - \bar{x}^*)^2}{B-1}$

$\bar{x}^* = \frac{\sum_{i=1}^B x_i^*}{B}$ and

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

Bias = statistic - parameter

Statistic - $E(\text{Statistic})$

$= 0$

$= S(x) - E[S(x)]$

$= S(x) - \hat{\lambda}(F)$

where $\hat{\lambda}(F)$ is the estimated value of the respective statistic, which is known as plug-in estimate.

H.W → How can you improve the bias by using Bootstrap sample.

Also find out 95% C.I using different Statistic including median, standard

Exam 2002

deviation, (mode) instead of mean and variance.

Solution

We will now find bootstrap 95% CI

$$X = \{4, 5, 6, 8, 9, 10, 10, 12\}$$

Sample no	Bootstrap sample	Median	Std	mode
1	{5, 5, 6, 8, 10, 10, 12, 9}	8.5	2.4	5, 10
2	{6, 6, 6, 8, 9, 10, 12, 12}	8.5	2.2	6
3	{4, 5, 8, 8, 9, 10, 10, 12}	8.5	2.5	8, 10
4	{5, 6, 6, 8, 9, 9, 10, 10}	8	1.7	6, 9, 10
5	{4, 4, 6, 8, 9, 10, 12, 12}	8.5	2.8	4, 12

medians:

$$\{8.5, 8.5, 8.5, 8, 8.5\}$$

$$\text{sorted} = \{8.0, 8.5, 8.5, 8.5, 8.5\}$$

$$CI = \{8, 8.5\}$$

For s.d

$$\{2.4, 2.2, 2.5, 1.7, 2.8\}$$

$$\text{sorted} = \{1.7, 2.2, 2.4, 2.5, 2.8\}$$

$$\text{C.I.} = (1.7, 2.8)$$

$$\text{mode CI} = (6, 10)$$

Lecture - 13 & 14

Vector of Parameters:

20-7-25

$$f(x_i, \theta)$$

$X_1, X_2, \dots, X_n$ are random variables with parameters $\theta_1, \theta_2, \dots, \theta_k$

$$\underline{\theta} = \{\theta_1, \theta_2, \dots, \theta_k\}$$

$T_1, T_2, \dots, T_k$ are estimators

$$E\{T_1, T_2, \dots, T_k\} = \{T_1, T_2, \dots, T_k\} = \{\theta_1, \theta_2, \dots, \theta_k\}$$

Assuming that $X_1, X_2, \dots, X_n$ be a random sample with size n then the density

function

$f(x; \theta_1, \theta_2, \dots, \theta_k)$ where

$\theta = (\theta_1, \theta_2, \dots, \theta_k)$ are the parameters

and $\bar{\theta} = \{\theta_1, \theta_2, \dots, \theta_k\}$ is the parameter space or vector of parameters

We are interested to estimate, $T_1(\theta) = \theta_1$,

$T_2(\theta) = \theta_2, \dots, T_k(\theta) = \theta_k$

If we let $k=1$, we get the point estimate. But a very special case is $k=n$. In this case, we will get n number of parameters. A point estimator of $T_1(\theta), T_2(\theta), \dots, T_n(\theta) = \theta_n$ is a vector of statistic as

$T_1, T_2, \dots, T_n$

$$T_1 = t_1(x_1, \dots, x_n) = \tau_1(\theta)$$

$$T_2 = t_2(x_1, \dots, x_n) = \tau_2(\theta)$$

$$T_n = t_n(x_1, \dots, x_n) = \tau_n(\theta)$$

where, $T_1, T_2, \dots, T_n$ are the n sample observations and $\tau_1(\theta), \tau_2(\theta), \dots, \tau_n(\theta)$ are the function of parameter θ .

Our next task is to find out how the distⁿ of the estimators $T_1, T_2, \dots, T_n$ be concentrated around $\tau_1(\theta), \tau_2(\theta), \dots, \tau_n(\theta)$.

We can find out the concentrations of the estimators $T_1, T_2, \dots, T_n$ to the parameters $\tau_1(\theta), \tau_2(\theta), \dots, \tau_n(\theta)$ by using the following four methods.

(a) Vector of variances.

(b) Linear combination of variances

(c) Ellipsoid of concentration.

(d) Wilk's generalised variance.

$[T, T]$

H.W
Ellipsoid



Ellipsoid of Concentrations;

The ellipsoid of concentration region around the mean that contains the most likely values of the vectors of estimators, $T = (T_1, T_2, \dots, T_n)$. Let

$T_1, T_2, \dots, T_n$ be an unbiased estimator

of the popⁿ parameters, $\tau_1(\theta), \tau_2(\theta), \dots$

and $\tau_n(\theta)$ be an unbiased estimator of

the popⁿ parameters $\tau_1(\theta), \tau_2(\theta), \dots$

$\tau_n(\theta)$. Also let $\sigma^{ij}(\theta)$ be the ij th

element of the inverse of the

covariance matrix of the estimator

$(T_1, T_2, \dots, T_n)$ where $\sigma^{ij}(\theta) =$

$$\text{cov}[T_i, T_j]$$

H.W $\rightarrow$ (do practical example of Ellipsoid of concentration and Wilk's generalised variance)

Person	Height	Weight
A	160	60
B	165	65
C	170	70
D	175	75

$\bar{x} = \begin{bmatrix} 167.5 \\ 67.5 \end{bmatrix}$

$$\text{var}(h) = \frac{(160 - 167.5)^2 + \dots + (175 - 167.5)^2}{n-1} = 41.67$$

$$\text{cov}(h, w) = \frac{(160 - 167.5)(60 - 67.5) + \dots + (175 - 167.5)(75 - 67.5)}{n-1}$$

$$= 41.67$$

$$S = \begin{bmatrix} 41.67 & 41.67 \\ 41.67 & 41.67 \end{bmatrix}$$

Exam 10/15/13

Lecture 15 & 16

27-07-25

Minimum Variance Bound Estimator

$$X_1, X_2, \dots, X_n \rightarrow f(X; \theta)$$

$$\theta_1, \theta_2, \dots, \theta_m$$

$$\text{var}(\theta_1) \leq \text{var}(\theta_i) \quad ; i = 2, \dots, m$$

↳ Minimum variance bound estimator

Book - Stuart Kendall

* Prove that,

the minimum variance Bound estimator
is an unique estimator.

~~Proof~~ In many situation when ~~at~~ although
the condition of attaining of MVB
estimator satisfy but still condition
there ~~many~~ may be a MV estimator
exists, it is unique, irrespective of
whether a MVB exists or not.

Let t_1 and t_2 be two distinct estimators. Also

let t_1 and t_2 be two MV unbiased estimators of $\tau(\theta)$, each with variance

let t_3 be a new estimator such that

$$t_3 = \frac{1}{2}(t_1 + t_2)$$

$$E(t_3) = \frac{1}{2} [E(t_1) + E(t_2)] = \frac{\tau(\theta) + \tau(\theta)}{2} = \tau(\theta)$$

$$\text{var}(t_3) = \text{var}\left[\frac{1}{2}(t_1 + t_2)\right]$$

$$= \frac{1}{4} \text{var}(t_1 + t_2)$$

$$= \frac{1}{4} \text{var}(t_1) + \frac{1}{4} \text{var}(t_2) + \frac{2 \text{cov}(t_1, t_2)}{4}$$

$$= \frac{1}{4} \cdot V + \frac{1}{4} V + 2 \text{cov}(t_1, t_2) \quad \text{--- ①}$$

According to the Cauchy-Schwarz inequality we get,

$$\text{cov}(t_1, t_2) = \iint [t_1 - \tau(\theta)] [t_2 - \tau(\theta)] L \, d\pi_1 d\pi_2 \dots d\pi_n$$

$$\leq \left[\int \dots \int \left[t_1 - \tau(\theta) \right]^2 d\mu_1, \mu_2, \dots, \mu_n \right]^{1/2}$$

$$\leq \left[\int \dots \int \left[t_2 - \tau(\theta) \right]^2 d\mu_1, \mu_2, \dots, \mu_n \right]^{1/2}$$

$$\leq \sqrt{\text{var}(t_1) \times \text{var}(t_2)}$$

$$\leq \sqrt{v \times v}$$

$$\leq v$$

$$\leq v$$

From eqⁿ (i) we get

$$\text{var}(t_3) \leq \frac{1}{4} (v+v)^2 + \frac{1}{4} \cdot 2v$$

the eqⁿ (ii) contradicts the assumption

that t_1 and t_2 have 0 MV unless

the equality holds

This means that the equality sign holds in eqⁿ (ii)

$$\text{cov}(t_1, t_2) = V \quad \text{--- (iii)}$$

We can also say that

$$[t_1 - \tau(\theta)] \propto [t_2 - \tau(\theta)]$$

$$\Rightarrow [t_1 - \tau(\theta)] = k(\theta) [t_2 - \tau(\theta)] \quad \text{--- (iv)}$$

$$\therefore \text{cov}(t_1, t_2) = k(\theta) \cdot \text{var}(t_2)$$

$$t_1 - \tau(\theta) = 1 [t_2 - \tau(\theta)]$$

$$\therefore t_1 = t_2$$

UMVUE:

Let $x_1, x_2, \dots, x_n$ be a random sample from

$f(\cdot; \theta)$. As estimator $T^* = f(t(x_1, \dots, x_n))$ of

$\tau(\theta)$ is defined to be an UMVUE of

$\tau(\theta)$ iff

$$\textcircled{i} =$$

Lower bound for variance:

1st tutorial syllabus:

1. Median and Model unbiased estimator
2. Jackknife and Bootstrap
3. Properties of good estimators
4. Cramer-Rao inequality, Bhattacharyya Bound.
5. Pitman estimator for location and scale.

Lower bound for variance:

Let, $X_1, X_2, \dots, X_n$ be a random sample drawn from $f(x_i; \theta)$, $\theta \in \bar{\Theta}$. Also let

$T = t(X_1, X_2, \dots, X_n)$ be an unbiased estimator of $\tau(\theta)$. Here we will consider that

$f(x_i; \theta)$ is a prob^y density f^n , is the analogous for discrete density f^n .

We make the following assumptions, called regularity conditions:

(i) $\frac{d}{d\theta} \log f(x_i; \theta); \theta \in \underline{\theta}$ and for all x exists

(ii) $\frac{d}{d\theta} \int \dots \int \prod_{i=1}^n f(x_i; \theta) dx_1, \dots, dx_n =$

$$\int \dots \int \frac{\partial}{\partial \theta} \prod_{i=1}^n f(x_i; \theta) dx_1, \dots, dx_n$$

(iii) $\frac{d}{d\theta} \int \dots \int t(x_1, x_2, \dots, x_n) \prod_{i=1}^n f(x_i; \theta) dx_1, \dots, dx_n$

$$= \int \dots \int t(x_1, x_2, \dots, x_n) \frac{\partial}{\partial \theta} \prod_{i=1}^n f(x_i; \theta) dx_1, \dots, dx_n$$

(iv) $0 < E_{\theta} \left[\frac{\partial}{\partial \theta} [\log f(x, \theta)]^2 \right] < \infty$ for all $\theta \in \underline{\theta}$

Prove that, the Cramér Rao inequality :

$$\text{Var}[T] \geq \frac{[\tau'(\theta)]^2}{n E_0 \left[\left[\frac{\partial}{\partial \theta} \log f(x; \theta) \right]^2 \right]}$$

Proof:

If the frequency function of discrete or continuous is $f(x; \theta)$, we may define the likelihood function of a sample of n independent samples by

$$L(x_1, x_2, \dots, x_n / \theta) = f(x_1 / \theta) f(x_2 / \theta) \dots f(x_n / \theta)$$

where, L is the joint frequency distribution of

$$\int \dots \int L dx_1 dx_2 \dots dx_n = 1 \quad \text{--- (1)}$$

Now differentiating both sides w.r. to θ we get,

$$\int \dots \int \frac{\partial L}{\partial \theta} dx_1, dx_2, \dots, dx_n = 0$$

which we may rewrite,

$$E\left(\frac{\log L}{2\theta}\right) = \int \dots \int \left(\frac{1}{L} \cdot \frac{\partial L}{\partial \theta}\right) L \cdot \frac{\partial x_1 \partial x_2 \dots \partial x_n}{\sqrt{J(\theta)}} = 0 \quad \text{--- (2)}$$

$$\Rightarrow \int \dots \int \left\{ \left(\frac{1}{L} \cdot \frac{\partial L}{\partial \theta}\right) \frac{\partial L}{\partial \theta} + L \cdot \frac{\partial}{\partial \theta} \left[\frac{1}{L} \cdot \frac{\partial L}{\partial \theta} \right] \right\} dx_1 \cdot dx_2 \dots dx_n = 0$$

$$\Rightarrow \int \dots \int \left\{ \left(\frac{1}{L} \cdot \frac{\partial L}{\partial \theta}\right)^2 + \frac{\partial \log L}{\partial \theta^2} \right\} L dx_1, dx_2, \dots, dx_n = 0$$

$$\Rightarrow \int \dots \int \left(\frac{1}{L} \cdot \frac{\partial L}{\partial \theta}\right)^2 L dx_1, dx_2, \dots, dx_n =$$

$$\int \dots \int \frac{\partial^2 \log L}{\partial \theta^2} \times L \cdot \frac{\partial x_1 \dots \partial x_n}{\sqrt{J(\theta)}}$$

from (2)

$$\Rightarrow E\left[\left(\frac{\log L}{2\theta}\right)^2\right] = -E\left[\frac{\partial^2 \log L}{\partial \theta^2}\right] \quad \text{--- (3)}$$

Now considering an unbiased estimator of t of $\tau(\theta)$, we have

$$E(t) = \int \dots \int t \cdot L dx_1, \dots, dx_n = \tau(\theta) \quad \text{--- (4)}$$

Taking 1st difference w.r.t. to θ ,

$$\int \dots \int t \cdot \frac{\partial L}{\partial \theta} dx_1, \dots, dx_n = \tau'(\theta)$$

$$\Rightarrow \int \dots \int t \left(\frac{1}{L} \cdot \frac{\partial L}{\partial \theta} \right) L dx_1, \dots, dx_n = \tau'(\theta)$$

$$\Rightarrow \int \dots \int t \frac{\partial \log L}{\partial \theta} L dx_1 dx_2 \dots dx_n = \tau'(\theta)$$

Now, applying Cauchy Schwartz inequality,

$$\{\tau'(\theta)\}^2 \leq \int \dots \int [t - \tau(\theta)]^2 L dx_1, \dots, dx_n \times$$

$$\int \dots \int \left(\frac{\partial \log L}{\partial \theta} \right)^2 L dx_1, \dots, dx_n$$

$$\Rightarrow [\tau'(\theta)]^2 \leq \text{var}(t) \times E \left(\frac{\partial \log L}{\partial \theta} \right)^2$$

$$\Rightarrow \text{var}(t) \geq \frac{[\tau'(\theta)]^2}{E \left(\frac{\partial \log L}{\partial \theta} \right)^2} \text{ which is}$$

Cramer - Rao Inequality

Lecture 19 & 20

04-08-25

Chapman Robbins and Kiefer inequality . .
for lower bound of variance

Cramér-Rao inequality

$$\text{var}(t) \geq \frac{\{t'(\theta)\}^2}{E\left[\frac{\partial \log L}{\partial \theta}\right]^2}$$

CRK:

CRK gives a lower bound for the
variance of an estimate but does not
require regularity conditions of the
Cramér-Rao inequality.

Theorem:

Let $X_1, X_2, \dots, X_n$ be a random
variables drawn from the
prob^y density $f^n f(x; \theta)$, $\theta \in \Theta$.

Also let $Y_1, Y_2, \dots, Y_n$ be the random
variables with prob^y density g^n

$f(x; \varphi)$, $\varphi \in \Theta$. Also let T be an unbiased estimator of both $\tau(\theta)$ and $\tau(\varphi)$ with $E[T^2] < \infty, \forall \theta \in \Theta$

Theorem:

If $\theta \neq \varphi$, assume that $f(x; \theta)$ and $f(x; \varphi)$ are different

Also, if $S(\theta) \supset S(\varphi)$

$$\Rightarrow [f(x; \theta) > 0] \supset [f_{\varphi}(x; \varphi) > 0]$$

Then

~~var~~ ~~var~~

$$\text{var}(T) \geq \sup \{ \varphi : S(\varphi) \subset S(\theta), \varphi \neq \theta \}$$

$$\frac{[\tau(\varphi) - \tau(\theta)]^2}{\text{var}(f(x; \varphi)/f(x; \theta))}$$

$$\text{var}(f(x; \varphi)/f(x; \theta))$$

$$\left\{ 1 - \frac{(\varphi, x)}{(\theta, x)} \right\} \text{var}$$

Proof:

Since T is an unbiased estimator for $\eta(\theta)$.

So, we may write, $E[T(x)] = \eta(\theta); \forall \theta \in \bar{\theta}$

Here, $\phi = \theta$,

$$\int_{S(\theta)} \frac{f(x, \phi) - f(x, \theta)}{f(x, \theta)} \times f(x, \theta) d\mu$$

$$= \eta(\phi) - \eta(\theta), \text{ which yield}$$

$$\text{cov} \left\{ T(x), \left[\frac{f(x, \phi)}{f(x, \theta)} - 1 \right] \right\} = \eta(\phi) - \eta(\theta)$$

Now, applying Cauchy-Schwartz inequality, we get

$$\text{cov}^2 \left\{ T(x), \left[\frac{f(x, \phi)}{f(x, \theta)} - 1 \right] \right\} \leq \text{var}(T(x)) \times \text{var} \left\{ \frac{f(x, \phi)}{f(x, \theta)} - 1 \right\}$$

$$\Rightarrow [\tau(\phi) - \tau(\theta)]^2 \leq \text{var}(\tau(x)) \times \text{var} \left\{ \frac{f(x; \phi)}{f(x; \theta)} - 1 \right\}$$

$$\Rightarrow \text{var}(\tau(x)) \geq \frac{[\tau(\phi) - \tau(\theta)]^2}{\text{var} \left\{ \frac{f(x; \phi)}{f(x; \theta)} - 1 \right\}}$$

* Bhattacharyya Bound:

We can find better lower bound than the minimum variance bound estimator, i.e.

$$\text{var}(t_n) \geq \frac{\{\tau'(\theta)\}^2}{E \left[\frac{\partial \log L}{\partial \theta} \right]^2} \quad \text{for the}$$

variance of an estimator. In this case, MVB is not attainable.

From the condition of MVB we have an estimator t_n for which $t_n - \tau(\theta)$ is a linear function of $\frac{\partial \log L}{\partial \theta} = \frac{1}{L} \cdot \frac{\partial L}{\partial \theta}$

But even if no such estimator exists, there still may be one for which $t_n - \tau(\theta)$ is a linear function of $\frac{\partial \log L}{\partial \theta}$

$$\frac{\partial^2 \log L}{\partial \theta^2}, \dots, \frac{\partial^n \log L}{\partial \theta^n}$$

Using (B4) we can write such a linear function that is closely related to $t_n - \tau(\theta)$

Let us consider t_n be an estimation of $\tau(\theta)$, we may write

$$L^{(n)} = \frac{\partial^n L}{\partial \theta^n} \quad \text{and} \quad \tau(\theta)^{(n)} = \frac{\partial^{(n)} \tau(\theta)}{\partial \theta^{(n)}}$$

$$D_S = t_n - Q(\theta) - \sum_{n=1}^{\infty} a_n \frac{L^n}{L}$$

where the constants to be determined

If we differentiate $\int \dots \int L \, d\mu_1 \, d\mu_2 \dots$

n times we get

$$\int \dots \int \frac{\partial L^n}{\partial L} L \, d\mu_1 \, d\mu_2 \dots d\mu_n = 0$$

$$\Rightarrow E\left[\frac{L^n}{L}\right] = \frac{1}{L} \int \dots \int L^n \, d\mu_1 \, d\mu_2 \dots d\mu_n$$

$$E[D_S] = E\left[t_n - Q(\theta) - \sum_{n=1}^{\infty} a_n \frac{L^n}{L}\right]$$

$$\therefore E[D_S^2] = \int \dots \int \left\{ t_n - Q(\theta) - \sum_{n=1}^{\infty} a_n \frac{L^n}{L} \right\}^2 \times L \, d\mu_1 \, d\mu_2 \dots d\mu_n$$

$$\int \dots \int \left[\frac{1}{L} \times \frac{1}{L} \right] L \, d\mu_1 \, d\mu_2 \dots d\mu_n$$

Now Diffⁿ p times we get

$$\int \dots \int \left[t_n - \tau(\theta) - \sum a_n \times \frac{L^{(n)}}{L} \right] \times \frac{L^{(p)}}{L} \times$$

$$L dx_1, dx_2, \dots, dx_n = 0$$

As summation and integration are

interchangeable, this gives

$$\int \dots \int \left[t_n - \tau(\theta) \right] \times \frac{L^{(p)}}{L} L dx_1, dx_2, \dots, dx_n =$$

$$\sum_{n=0}^{\infty} a_n \cdot \frac{L^{(n)}}{L} \times \frac{L^{(p)}}{L} \times L dx_1, dx_2, \dots, dx_n$$

S.H.S

$$\times \left\{ \int \dots \int \left[t_n - \tau(\theta) \right] \times \frac{L^{(p)}}{L} \times L dx_1, dx_2, \dots, dx_n \right.$$

or $\dots \tau(\theta)$

$$R.H.S. = \sum_{n=0}^{\infty} a_n \int \dots \int \left[\frac{L^n}{L} \times \frac{L^{(p)}}{L} \right] L dx_1, dx_2, \dots, dx_n$$

$$= \sum_{n=1}^{\infty} a_n E \left[\frac{L^n}{L} \times \frac{L^p}{L} \right]$$

$$= \sum_{n=1}^{\infty} a_n \text{cov} \left[\frac{L^n}{L} \times \frac{L^p}{L} \right]$$

$$= \sum_{n=1}^{\infty} a_n J_{np}$$

$$\therefore \tau^{(p)}(\theta) = \sum_{n=1}^{\infty} a_n \cdot J_{np}$$

$$2) a_n = \sum_{p=1}^{\infty} \tau^{(p)}(\theta) J_{np}^{-1}$$

Now setting the value of a_n

in eqⁿ (1)

$$DS = f_n - \tau(\theta) - \sum_{n=1}^{\infty} \sum_{p=1}^{\infty} \tau^{(p)}(\theta) J_{np}^{-1} \left[\frac{L^{n'}}{L} \right] \frac{L}{L}$$

$$L dn_1, dn_2, \dots, dn_n$$

WVVV

1, 2, 4 Chapter

1st tutorial
chapter

Lecture - 29 & 32

6-8-25

Chapter-1 : Median and Modal unbiased

estimates.

Chapter-2: Jackknife and Bootstrap

Chapter-3: Cramer Rao Inequality,

Bhattacharyya lower bound

⑧ Invariance property:

estimator $\pm, \times, \div$ corr

Location

Location invariant:

As estimator $T = t(x_1, x_2, \dots, x_n)$ is

defined to be location invariant if

$$t(x_1 + c, x_2 + c, \dots, x_n + c) = t(x_1, x_2, \dots, x_n) + c$$

for all values $x_1, x_2, \dots, x_n$ and all c

Prove that,

① $\frac{Y_1 + Y_n}{2}$ is a location invariant.

② Prove that S^2 is not a location invariant estimator.

③ Prove that $\frac{Y_n - Y_1}{2}$ is not a location invariant estimator.

④ Give a practical example of location invariant estimator.

⑤ $\bar{X}$ is a location invariant estimator.

⑥ According to the definition,

$$f(x_1 + c, x_2 + c, \dots, x_n + c) = \bar{x}_n$$

$$\bar{x}_n = \frac{x_1 + x_2 + \dots + x_n}{n}$$

$$= \frac{(x_1 + c) + (x_2 + c) + \dots + (x_n + c)}{n}$$

$$= \frac{x_1 + x_2 + \dots + x_n}{n} + c$$

$$= \bar{x}_n + c$$

Scale Invariant:

An estimator $T = t(x_1, x_2, \dots, x_n)$ is defined to be scale invariant iff $t(cx_1, cx_2, \dots, cx_n) = c t(x_1, x_2, \dots, x_n)$ for all values of $x_1, x_2, \dots, x_n$ for all values of c

HW Prove that, $\bar{x}_n, \sqrt{s^2}, \frac{Y_1 + Y_n}{2}$ and

$Y_n - Y_1$ are scale invariant.

⊗ Location parameters:

Let $\{f(x; \theta) ; \theta \in \bar{\Theta}\}$ be a family of densities p with parameter θ . The parameter θ is defined to be a

location parameter if and only if the density $f(x; \theta)$ can be written as a f^n of $x - \theta$ that is $f(x; \theta) = h(x - \theta)$ for some f^n of $h(\cdot)$ equivalently.

Pitman estimator for location:

Let $x_1, x_2, \dots, x_n$ denote a random sample drawn from the density f^n $f(x; \theta)$, where θ is a location parameter and $\theta \in \underline{\Theta}$

Then the pitman estimator is written by

$$\hat{\theta}(x_1, x_2, \dots, x_n) = \frac{\int_{\underline{\Theta}} \theta \prod_{i=1}^n f(x_i; \theta) d\theta}{\int_{\underline{\Theta}} \prod_{i=1}^n f(x_i; \theta) d\theta} \quad \text{--- (1)}$$

The estimator in eqⁿ (1) is called pitman estimator for location parameter, which has uniformly smallest mean squared error within the class of location-invariant estimators.

Examples

37, 38, 39

chapter - 9

mod 98 12

Page (333, 334, 375)

Scale parameter:

Let $\{f(x; \theta), \theta \in \bar{\Theta}\}$ be a family of density densities with parameter θ . The parameter θ is defined to be scale parameter iff the density $f^n(x; \theta) = \frac{1}{\theta} h(x/\theta)$ for some density f^n & $h(\cdot)$.

Equivalently, θ is called scale parameter of the density $f^n(x; \theta)$ and the distⁿ of $f(x/\theta)$ does not depend on θ .

Pitman estimator for scale:

Let $x_1, x_2, \dots, x_n$ be a random sample drawn from the density $f^n f(x; \theta); \theta > 0$ is a

scale parameter. Assume that $f(x; \theta) = 0$ for $x \leq 0$ that means the random variable x_i only takes positive values. within the class of scale-invariant estimators, the pitman estimator

for scale is

$$t(x_1, x_2, \dots, x_n) = \frac{\int_0^\infty \frac{1}{\theta^n} \prod_{i=1}^n f(x_i; \theta) d\theta}{\int_0^\infty \frac{1}{\theta^3} \prod_{i=1}^n f(x_i; \theta) d\theta}$$

which has smallest risk the loss L^n

$$L(t; \theta) = (t - \theta)^2 / \theta$$

H.W: (ex (41), (42))

(C374) (42)

⊗ When we apply Robust Estimator?

In practice, we make many assumptions in parametric inference, and any one or all of this may seldom have

accurate knowledge about the underlying distⁿ. Similarly the assumption of

event it may not hold. Any

test or estimator that perform well

under model modification or under

assumption is usually referenced as

robust estimator.

⊗ How can you construct the robust estimator?

Let $X_1, X_2, \dots, X_n$ be a random

sample drawn from $I_0 \sim N(\mu, \sigma^2)$

with 100% and from $I_1 \sim N(\mu, \sigma^2)$

$k > 1$ with $(1-\alpha)100\%$ of the time where both N and σ^2 is unknown. We are interested to estimate μ .

We may assume that really the sample is drawn from,

$f(x) = \alpha f_0(x) + (1-\alpha)f_1(x)$; clearly $\bar{x}$ is still unbiased estimator of μ

Still unbiased estimator

If α is nearly 1 there is no problem, since the underlying distⁿ is nearly $N(\mu, \sigma^2)$ and $\bar{x}$ is uniformly minimum variance unbiased estimator of

μ with $V(\bar{x}_n) = \frac{\sigma^2}{n}$

If $(1-\alpha)$ large, then we are sampling α from f then the variance of $v(x)$ is σ^2 with prob^y α and $k^2\sigma^2$ with prob^y $(1-\alpha)$ and we have,

$$\begin{aligned} V(\bar{x}) &= \frac{\text{var}(x)}{n} \\ &= \frac{\alpha\sigma^2}{n} + \frac{(1-\alpha)k\sigma^2}{n} \\ &= \frac{\sigma^2}{n} [\alpha + (1-\alpha)k] \end{aligned}$$

for heterogeneous characteristic

But, if it is homogeneous, $\alpha = 1$

If $\alpha > 1$, then $v(\bar{x})$ will be more large

If $k(1-\alpha)$ is large, $\text{var}(\bar{x})$ will also be large.

Even for a single wild obsⁿ makes $\bar{x}$ a considerable error

It is known that sample median is more better estimator than mean in process of extreme values

If $Z_{1/2} = M$ be the sample median drawn from $X_1, \dots, X_n$ as an estimate of central value μ , which is known as population median.

Then for large n , it is true that

$$E (m - \mu)^2 = \text{var} (m)$$

$$\approx \frac{1}{4n} \times \frac{4n [f(\mu)^2]}{\pi \sigma^2} = \frac{\pi \sigma^2}{2n}$$

where $f(\mu)$ is the distⁿ of the popⁿ median.

Since $f(\mu) = \alpha f_0(\mu) + (1-\alpha)f_1(\mu)$

$$f(\mu) = \alpha \cdot \frac{1}{\sigma\sqrt{2\pi}} + (1-\alpha) \cdot \frac{1}{\sigma\sqrt{2\pi k}}$$

$$= \frac{1}{\sigma\sqrt{2\pi}} \left(\alpha + \frac{1-\alpha}{\sqrt{k}} \right)$$

$$\therefore E(m-\mu)^2 = \frac{1}{4n \left\{ \left[\alpha + \frac{1-\alpha}{\sqrt{k}} \right]^2 \cdot \frac{1}{\sigma^2} \right\}}$$

$$\Rightarrow \text{var}(m) = \frac{\pi\sigma^2}{2n} \cdot \frac{1}{\left(\alpha + \frac{1-\alpha}{\sqrt{k}} \right)^2}$$

$$(m) \text{ var } = \frac{1}{n} \quad (2)$$

If $k \rightarrow \alpha$, then the eqⁿ (2) becomes

$$\text{var}(m) = \frac{\pi\sigma^2}{2n} \times \frac{1}{\alpha^2}$$

If $\alpha = 1$, $\text{var}(m) = \frac{\pi\sigma^2}{2n}$ that means

all the observations are taking from second pdf i.e., homogeneity will be

high i.e., $\text{var}(m)$ is not affected by k

This indicates that the estimator m will not be greatly affected by how large k is.

$$\begin{aligned}\therefore \frac{\text{var}(\bar{x})}{\text{var}(m)} &= \frac{\frac{\sigma^2}{n} [\alpha + (1-\alpha)k]}{\frac{\pi \sigma^2}{2n} \left[\frac{1}{\alpha + \frac{1-\alpha}{\sqrt{k}}} \right]^2} \\ &= \frac{\cancel{\sigma^2}}{\cancel{\pi}} [\alpha + (1-\alpha)k] \times \frac{2n \left\{ \alpha + \frac{1-\alpha}{\sqrt{k}} \right\}^2}{\cancel{\pi \sigma^2}} \\ &= \frac{2}{\pi} [\alpha + (1-\alpha)k] \left[\alpha + \frac{1-\alpha}{\sqrt{k}} \right]^2\end{aligned}$$

If $k \rightarrow \infty$

$$\therefore \frac{\text{var}(\bar{x})}{\text{var}(m)} = \alpha$$

Where as $\text{var}(\bar{x}) \rightarrow \infty$ as $k \rightarrow \infty$

$$\therefore \text{var}(m) = \frac{\pi \sigma^2}{2n \alpha^2}$$

So we may write

$$\therefore \text{var}(m) < \text{var}(\bar{x})$$

So, the estimator m is called a robust estimator