

Statistical Signal Processing
Prof. Prabin Kumar Bora
Department of Electronics & Electrical Engineering
Indian Institute of Technology - Guwahati

Lecture 9

Estimation Theory 2: MVUE and Cramer Rao Lower bound

(Refer Slide Time: 00:50)

Let us recall

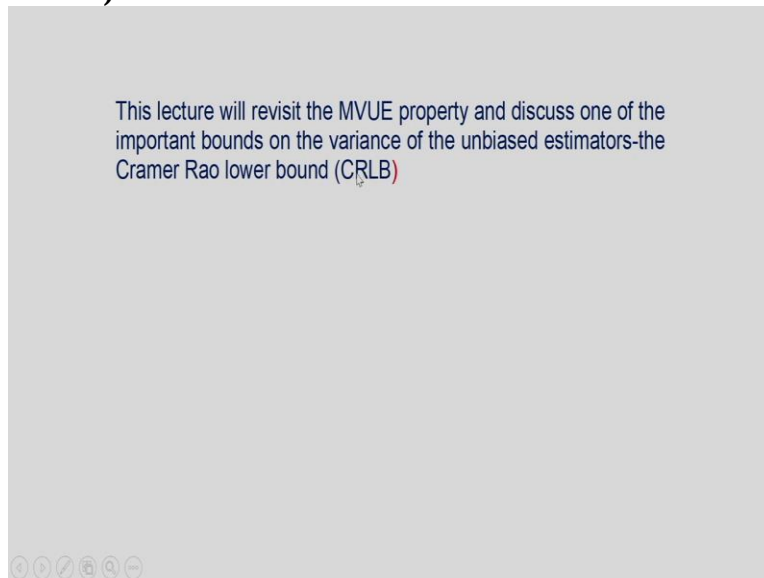
- ❖ The observed random data $X_1, X_2, \dots, X_N$ are characterized by a joint PDF which depends on some unobservable parameter θ . An estimator $\hat{\theta}(X) = \hat{\theta}(X_1, X_2, \dots, X_N)$ is a function by which we guess about the value of the unknown parameter θ .
- ❖ An estimator $\hat{\theta}$ of θ is unbiased if and only if $E\hat{\theta} = \theta$.
- ❖ A minimum variance unbiased estimator (MVUE) has the lowest variance in the class of unbiased estimators.
- ❖ $\hat{\theta}$ is called consistent if
$$\lim_{N \rightarrow \infty} P(|\hat{\theta} - \theta| \geq \varepsilon) = 0 \text{ for any } \varepsilon > 0$$
If $\hat{\theta}$ is unbiased and $\lim_{N \rightarrow \infty} \text{var}(\hat{\theta}) = 0$, then $\hat{\theta}$ is consistent.
- ❖ The mean square error (MSE) of an estimator is given by
$$MSE = E(\theta - \hat{\theta})^2$$
and minimizing the MSE is an important estimation criterion.

Hello students, welcome to lecture 10 on Estimation theory 2 it will cover MVUE and Cramer Rao lower bound. Let us recall the observed random data X_1, X_2 up to X_N are characterized by a joint PDF which depends on some unobserved parameter θ and estimator $\hat{\theta}(X)$ is a function by which we guessed about the value of the unknown parameter θ . So, corresponding to suppose random data we have a PDF, which is characterized by a parameter θ . And estimator $\hat{\theta}(X)$ is a function by which we guessed about the value of the unknown parameter θ .

An estimator $\hat{\theta}$ of parameter θ is unbiased if $E\hat{\theta} = \theta$. So, this is the requirement for an unbiased estimator. A minimum variance unbiased estimator MVUE has the lowest variance in the class of unbiased estimators. So this is 1 of the important estimator that is minimum variance unbiased estimator MVUE and $\hat{\theta}$ that is estimator $\hat{\theta}$ is called consistent if this probability of ε arbitrary there is summation goes down to 0 as N tends to infinity.

So, this is the definition of consistent estimator. Eventually as we have more and more data, θ as will grow more and more close to θ that is the interpretation of consistent estimator. If θ had this unbiased and limit of variance of θ hat = 0 then θ is consistent that is a simple test for consistent estimator when θ is unbiased. The mean square error MSE of an estimator is given by this $MSE = E$ of $\theta - \theta$ hat whole square and minimizing the MSE is an important estimation criteria.

(Refer Slide Time: 03:13)



This lecture will revisit the MVUE property and discuss 1 of the important bounds on the variance of the unbiased estimators the Cramer Rao lower bond it is abbreviated as CRLB.

(Refer Slide Time: 03:33)

Minimum Variance Unbiased Estimator

❖ Recall that $\text{var}(\theta) = E(\theta - \hat{\theta})^2$

❖ $\hat{\theta}$ is an MVUE if

$$E(\hat{\theta}) = \theta$$

and $\text{Var}(\hat{\theta}) \leq \text{Var}(\hat{\theta}')$

where $\hat{\theta}'$ is any other unbiased estimator of θ .

❖ MVUE is also known as the best unbiased estimator

Minimum Variance unbiased estimator that is MVUE recall that variance of theta as it define it as E of theta - theta hat whole squared. We can write it as theta hat - theta, but it does not matter. I mean both where we can write theta hat is an MVUE if E of theta hat = theta and variance of theta hat is less than equal to variants of theta hat depth where theta hat depth is any other unbiased estimator of theta MVUE also known as the best unbiased estimator because its variance is lowest. Therefore it is also known as the best unbiased estimator.

(Refer Slide Time: 04:24)

MVUE uniqueness theorem: An MVUE is unique

❖ Suppose $\hat{\theta}_1$ and $\hat{\theta}_2$ are two MVUEs for the deterministic parameter θ .

❖ Clearly, $E\hat{\theta}_1 = E\hat{\theta}_2 = \theta$. Suppose $\text{Var}(\hat{\theta}_1) = \text{Var}(\hat{\theta}_2) = \sigma^2$

❖ Assume another estimator $\hat{\theta}_3 = \frac{\hat{\theta}_1 + \hat{\theta}_2}{2}$

$$\begin{aligned} \text{var}(\hat{\theta}_3) &= \frac{\text{var}(\hat{\theta}_1) + \text{var}(\hat{\theta}_2) + 2\text{cov}(\hat{\theta}_1, \hat{\theta}_2)}{4} \\ &\leq \frac{\text{var}(\hat{\theta}_1) + \text{var}(\hat{\theta}_2) + 2|\text{cov}(\hat{\theta}_1, \hat{\theta}_2)|}{4} \\ &\leq \frac{\text{var}(\hat{\theta}_1) + \text{var}(\hat{\theta}_2) + 2\sqrt{\text{var}(\hat{\theta}_1)\text{var}(\hat{\theta}_2)}}{4} \quad (\text{using C.S. inequality}) \\ &= \frac{\sigma^2 + \sigma^2 + 2\sigma^2}{4} \end{aligned}$$

$$\therefore \text{var}(\hat{\theta}_3) = \sigma^2 \Rightarrow \text{cov}(\hat{\theta}_1, \hat{\theta}_2) = \sigma^2$$

$$\begin{aligned} \text{Now, } \text{var}(\hat{\theta}_1 - \hat{\theta}_2) &= \text{var}(\hat{\theta}_1) + \text{var}(\hat{\theta}_2) - 2\text{cov}(\hat{\theta}_1, \hat{\theta}_2) \\ &= \sigma^2 + \sigma^2 - 2\sigma^2 = 0 \end{aligned}$$

$$\therefore \hat{\theta}_1 = \hat{\theta}_2 \quad \text{with probability 1}$$

We will prove 1 important theorem and MVUE is unique. Suppose theta 1 hat and theta 2 hat are 2 MVUE for the deterministic parameter theta clearly E of theta 1 hat = E of theta 2 hat = theta because of unbiased. Suppose variance of theta 1 hat = variance of theta 2 hat that = sigma

squared because both have the same minimum variance. So Let that minimum variance the σ^2 assume another estimator $\hat{\theta}_3$ which is given by $\hat{\theta}_1 + \hat{\theta}_2$ divided by 2, then variance of $\hat{\theta}_3$ will be = variance of $\hat{\theta}_1$ + variance of $\hat{\theta}_2$ hat + 2 times covariance of $\hat{\theta}_1$ $\hat{\theta}_2$ whole term divided by 4.

Now, this expression is less than equal to variance of $\hat{\theta}_1$ + variance of $\hat{\theta}_2$ + 2 times mod of covariance of $\hat{\theta}_1$ $\hat{\theta}_2$ divided by 4 because in a number is less than equal to its mod value. Now we will apply the Cauchy-Schwarz inequality, then our expression this expression will be less than equal to variance of $\hat{\theta}_1$ variance of $\hat{\theta}_2$ + 2 times square root of variance of $\hat{\theta}_1$ into variance of $\hat{\theta}_2$ hat.

Because the magnitude of debt covariance of $\hat{\theta}_1$ $\hat{\theta}_2$ is less than equal to square root of variance of $\hat{\theta}_1$ into variance of $\hat{\theta}_2$ hat this expression is now = σ^2 + σ^2 + this is σ^2 therefore + 2 σ^2 divided by 4, which = σ^2 . Therefore, what we have established at variance of $\hat{\theta}_3$ hat is less than equal to σ^2 , but it cannot be less than σ^2 .

Because minimum variance is σ^2 , therefore variance of $\hat{\theta}_3$ hat is also = σ^2 . And if I substitute this value in this expression, then we will get that covariance of $\hat{\theta}_1$ $\hat{\theta}_2$ is also = σ^2 . So, if there is another estimated $\hat{\theta}_3$ hat, which = $\hat{\theta}_1$ + $\hat{\theta}_2$ divided by 2, then its variance is σ^2 . And this source that covariance of $\hat{\theta}_1$ $\hat{\theta}_2$ must be = σ^2 .

And now, let us examine these expression variance of $\hat{\theta}_1$ – $\hat{\theta}_2$ hat this will be = variance of $\hat{\theta}_1$ + variance of $\hat{\theta}_2$ hat - 2 times covariance of $\hat{\theta}_1$ $\hat{\theta}_2$ hat. Now, subsidiary the value will get this is σ^2 , this is σ^2 and this is σ^2 . Therefore this expression will be = 0 that is variance of $\hat{\theta}_1$ - $\hat{\theta}_2$ hat = 0. This implies that $\hat{\theta}_1$ must be = $\hat{\theta}_2$ hat with probability 1. Therefore an MVUE is unique, minimum variance unbiased estimator is always unique. The uniqueness of MVUE makes it the most desirable estimator.

(Refer Slide Time: 08:15)

Cramer Rao lower bound (CRLB)

- ❖ The joint PDF or the joint PMF of the random vector $\mathbf{X} = [X_1 \ X_2 \ \dots X_N]^T$ determines the minimum variance.
- ❖ Let $f(\mathbf{x}; \theta) = f(x_1, \dots, x_N; \theta)$ be the joint PDF which characterizes $\mathbf{x}$. This function is also called *likelihood function*.
- ❖ $L(\mathbf{x}; \theta) = \ln f(x_1, x_2, \dots, x_N; \theta)$ is called the *log-likelihood function*.
- ❖ $I(\theta) = E\left(\frac{\partial L}{\partial \theta}\right)^2$ is a measure of average information of the random variables and is called *Fisher information statistic*.
- ❖ Cramer Rao theorem gives a lower bound of the variance of an unbiased estimator $\hat{\theta}$ in terms of $I(\theta)$.

Now we will discuss Cramer Rao lower bound CRLB the joint PDF or the joint PMF of the random vector $\mathbf{X}$ that is $X_1 \ X_2 \ \text{up to} \ X_N$ transpose determines the minimum variance the joint PDF will only determine the minimum variance of that $\hat{\theta}$ is a function of this random this random variables. Let $f(\mathbf{x}; \theta)$ that = small f PDF $X_1 \ X_2 \ \text{up to} \ X_N$; θ that is a it is a function of θ be the joint PDF which characterizes the random data vector $\mathbf{x}$.

This function is also called likelihood function. So, this is the likelihood function, which is a function of θ and we denoted by this is the likelihood function. Now, if we take the logarithm capital $L(\mathbf{x}; \theta)$ that is the log of this likelihood function is called a log likelihood function. Now, we define a term known as Fisher information statistics $I(\theta)$ is the expected value of $\frac{\partial L}{\partial \theta}$ whole square.

So, we take the partial derivative of L with respect to θ and then take square then they variance value of that quantity is a measure of average information and it is called Fisher information statistic like any information, whether it is also positive and we can solve that it is an increasing function Cramer Rao theorem gives a lower bound of the variance of an unbiased estimator $\hat{\theta}$ in terms of $I(\theta)$.

So, the lower and the lower bound in the variance will be given in terms of $I(\theta)$. So, we have discussed what is $I(\theta)$ that is the information Fisher information statistics that is E of $\frac{\partial L}{\partial \theta}$ squared.

(Refer Slide Time: 10:38)

Cramer Rao theorem

- ❖ Suppose $\hat{\theta}$ is an unbiased estimator of θ and $f(x_1, \dots, x_N; \theta)$ satisfies the following regularity conditions:
 - (i). θ lies in open interval D_θ .
 - (ii). The support $\chi = \{(x_1, \dots, x_N) \mid f(x_1, \dots, x_N; \theta) > 0\}$ does not depend on θ .
 - (iii). $\frac{\partial L}{\partial \theta}$ exists and is finite.
- ❖ Under the above regularity conditions, $Var(\hat{\theta}) \geq \frac{1}{I(\theta)}$
- ❖ The equality of CR bound holds if $\frac{\partial L}{\partial \theta} = c(\hat{\theta} - \theta)$ where c is a constant.

Suppose, $\hat{\theta}$ is an unbiased estimator of θ and $f(x_1, x_2, \dots, x_N; \theta)$ satisfies the following regularity condition likelihood function satisfies the following regularity conditions. That is θ lies in open interval D_θ there is an interval D_θ on which θ is defined. The support χ , this is defined by $x_1, x_2, \dots, x_N$ such that this likelihood function is greater than 0 wherever likelihood function is 0, that support does not depend on θ .

So essentially we will be using an integration that limited integration does not depend on θ $\frac{\partial L}{\partial \theta}$ partial derivative of L with respect to θ exist and is finite so these 3 are the regularity conditions under the above regularity conditions, variance of $\hat{\theta}$ is greater than equal to $\frac{1}{I(\theta)}$ where $I(\theta)$ is the Fisher information statistic the equality of CR bound, this is the Cramer Rao bound the equality of CR bound holds if $\frac{\partial L}{\partial \theta} = \text{scalar multiple of } \hat{\theta} - \theta$ where c is a constant scalar it may be a function of θ , but it does not involve $\hat{\theta}$.

So, these regularity conditions are purely general it is not very restrictive and under this regularity conditions Cramer Rao theorem is variance of $\hat{\theta}$ is greater than equal to $\frac{1}{I(\theta)}$

theta the equality holds when the partial derivative of log likelihood function can be written in this policy into theta hat - theta.

(Refer Slide Time: 12:51)

Cramer Rao theorem

- ❖ The proof of CR theorem is based on simple calculus and the Cauchy Schwartz (CS) inequality:

$$|\langle \mathbf{a}, \mathbf{b} \rangle|^2 \leq \|\mathbf{a}\|^2 \|\mathbf{b}\|^2$$

where the equality holds when $\mathbf{a} = c\mathbf{b}$, where c is any scalar.

- ❖ In the case of square-integrable functions on the real line, we can apply the CS inequality to get

$$\left(\int_{-\infty}^{\infty} f_1(x)f_2(x)dx \right)^2 \leq \int_{-\infty}^{\infty} f_1^2(x)dx \int_{-\infty}^{\infty} f_2^2(x)dx$$

with the equality holding when $f_1(x) = cf_2(x)$

We will try to prove this theorem. The proof of CR theorem is based on simple calculus and the Cauchy Schwartz inequality that is a very important inequality, we have discussed about this earlier and this inequality is given by in norm product this is in norm product magnitude of the in norm product square is less than equal to norm of a squared into norm of b squared, where equality holds when $\mathbf{a} = c$ times $\mathbf{b}$ where c is any scalar.

So, this is the statement of Cauchy Schwartz inequality in the case of square integrable functions on the real line, we can apply this CS inequality Cauchy Schwartz inequality to get this relationship that is integration from - infinity to infinity $f_1(x)$ into $f_2(x) dx$ whole squared. This was relation = integration - infinity to infinity $f_1^2(x) dx$ integration - infinity to infinity $f_2^2(x) dx$. So, this is the Cauchy Schwartz inequality in terms of 2 functions.

So, here this in norm product is defined in terms of this integration and norm is defined by this integration. So, that way this is the result we get and this result will be using for deriving the processor inequality also in the case of 2 functions, the equality that is, this quantity will be equal to product of this only when 1 function is a scalar multiple of the other function.

(Refer Slide Time: 14:48)

Proof:

❖ $\hat{\theta}$ is an unbiased estimator of θ , $\int_{-\infty}^{\infty} (\hat{\theta} - \theta) f(\mathbf{x}; \theta) d\mathbf{x} = 0$. $E(\hat{\theta} - \theta) = 0$

$$\Rightarrow \frac{\partial}{\partial \theta} \int_{-\infty}^{\infty} (\hat{\theta} - \theta) f(\mathbf{x}; \theta) d\mathbf{x} = 0.$$

❖ Since the limits of integration do not involve θ , we can write

$$\int_{-\infty}^{\infty} \frac{\partial}{\partial \theta} ((\hat{\theta} - \theta) f(\mathbf{x}; \theta)) d\mathbf{x} = 0.$$

$$\Rightarrow \int_{-\infty}^{\infty} (\hat{\theta} - \theta) \frac{\partial}{\partial \theta} f(\mathbf{x}; \theta) d\mathbf{x} - \int_{-\infty}^{\infty} f(\mathbf{x}; \theta) d\mathbf{x} = 0.$$

$$\therefore \int_{-\infty}^{\infty} (\hat{\theta} - \theta) \frac{\partial}{\partial \theta} f(\mathbf{x}; \theta) d\mathbf{x} = \int_{-\infty}^{\infty} f(\mathbf{x}; \theta) d\mathbf{x} = 1. \quad (1)$$

Now will prove this theorem $\hat{\theta}$ is an unbiased estimator of θ . This expression actually this is E of $\hat{\theta} - \theta$ must be $= 0$ so that if we write in terms of the definition. So we will get integration - infinity to infinity $\hat{\theta} - \theta$ into f of $\mathbf{x}$ θ $d\mathbf{x} = 0$, this $d\mathbf{x}$ is a vector so this integral actually this is a multiple integral. So we have supposed and variables x_1, x_2 up to x_N then this is an N holds integration.

Now, since the limits of integration do not involve θ , we can take this partial derivative inside here it was outside the integration now we can take it inside $\frac{\partial}{\partial \theta}$ of $\hat{\theta} - \theta$ into f $\mathbf{x}$ θ $d\mathbf{x}$, that must be $= 0$. This $\mathbf{x}$ is a factorial now since this is a product of 2 functions, we can apply the product rule and we will get integration - infinity to infinity $\hat{\theta} - \theta$ into partial derivative of f $\mathbf{x}$ θ with respect to θ into $d\mathbf{x}$.

Now $\hat{\theta} - \theta$ will have a partial derivative - 1 that way - integration - infinity to infinity, f of $\mathbf{x}$ θ $d\mathbf{x}$ must be $= 0$. So this is thus applying differentiation we get this. So we will take this quantity right hand side. So that way it will become positive and then also I know that this integration $= 1$. Therefore, this quantity integration $\hat{\theta} - \theta$ into $\frac{\partial}{\partial \theta}$ of f of $\mathbf{x}$ θ $d\mathbf{x}$ that must be $= 1$ this result is very important for us.

(Refer Slide Time: 17:06)

Proof:

❖ We have established $\int_{-\infty}^{\infty} (\hat{\theta} - \theta) \frac{\partial}{\partial \theta} f(\mathbf{x}; \theta) d\mathbf{x} = 1 \quad (1)$

❖ Note also that $\frac{\partial}{\partial \theta} f(\mathbf{x}; \theta) = \frac{\partial}{\partial \theta} (\ln f(\mathbf{x}; \theta)) f(\mathbf{x}; \theta) = \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} f(\mathbf{x}; \theta)$.

❖ Therefore, from (1) $\int_{-\infty}^{\infty} (\hat{\theta} - \theta) \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} f(\mathbf{x}; \theta) d\mathbf{x} = 1$

❖ Squaring both sides

$$\left(\int_{-\infty}^{\infty} (\hat{\theta} - \theta) \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} f(\mathbf{x}; \theta) d\mathbf{x} \right)^2 = 1$$

❖ We can rewrite the above as

$$\left(\int_{-\infty}^{\infty} (\hat{\theta} - \theta) \sqrt{f(\mathbf{x}; \theta)} \frac{\partial}{\partial \theta} L(\mathbf{x}; \theta) \sqrt{f(\mathbf{x}; \theta)} d\mathbf{x} \right)^2 = 1.$$

❖ Applying CS inequality to the left-hand side,

$$\left(\int_{-\infty}^{\infty} (\hat{\theta} - \theta) \sqrt{f(\mathbf{x}; \theta)} \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} \sqrt{f(\mathbf{x}; \theta)} d\mathbf{x} \right)^2 \leq \int_{-\infty}^{\infty} (\hat{\theta} - \theta)^2 f(\mathbf{x}; \theta) d\mathbf{x} \int_{-\infty}^{\infty} \left(\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} \right)^2 f(\mathbf{x}; \theta) d\mathbf{x}$$

$$\therefore 1 \leq \text{var}(\hat{\theta}) I_{\theta}(\theta)$$

Now, we have established that integration - infinity to infinity $\hat{\theta} - \theta \frac{\partial}{\partial \theta} f$ of $\mathbf{x}$ $\theta d\mathbf{x}$ that $= 1$ this is supposed to call this expression as 1 note that this partial derivative $\frac{\partial}{\partial \theta} f$ $\mathbf{x}$ θ that is equal to we can write in terms of log likelihood function like this log of $\frac{\partial}{\partial \theta} f$ of $\mathbf{x}$ θ into f $\mathbf{x}$ θ . So, this will be same as this quantity because 1 by f $\mathbf{x}$ θ will come and that and this will get cancelled.

So, that way this expression can be rewritten as this so that $\frac{\partial}{\partial \theta} f$ $\mathbf{x}$ θ is same as $\frac{\partial L}{\partial \theta}$ $\mathbf{x}$ θ $\frac{\partial}{\partial \theta} f$ $\mathbf{x}$ θ . So this is 1 important derivation. Therefore, this expression 1 can read as now integration - infinity to infinity $\hat{\theta} - \theta$ into $\frac{\partial L}{\partial \theta}$ $\mathbf{x}$ θ $\frac{\partial}{\partial \theta} f$ $\mathbf{x}$ θ $d\mathbf{x} = 1$. So, this is the modified form of this equation our aim is to apply the Cauchy Schwartz inequality to the left side of this expression.

So, that way we Cauchy Schwartz this expression if I squared then also this will be $= 1$. Now, we have to write it as a product of 2 functions. And since this quantity is greater than equal to 0. We can take this square root and that way we can write it in terms of 2 function 1 is $\hat{\theta} - \theta$ into root over f $\mathbf{x}$ θ and other 1 is that partial derivative of L $\mathbf{x}$ θ into square root of; f $\mathbf{x}$ $\theta d\mathbf{x}$, so that way we have 2 functions.

So this expression is rewritten as integration - infinity to infinity integration from - infinity to infinity of $\hat{\theta} - \theta$ in to root over f $\mathbf{x}$ θ into $\frac{\partial L}{\partial \theta}$ $\mathbf{x}$ θ into root over f $\mathbf{x}$

theta dx whole squared = 1. So, we have a known to product up to functions suppose if I call this as a, this is supposed this = a vector and this is b vector, this part is b vector, this part is a vector and this part is b vector, then we can apply this CS inequality.

So, therefore, this is the inner product of 2 functions squared. So, that way this squared must be less than equal to if I consider the individual function norm of individual function, this will be integration - infinity to infinity theta hat - theta whole squared, this is theta hat – theta whole squared and this is square root of f x theta therefore, squaring will get f of x theta dx. So, this is part and first norm square and norm of the second function square will be integration from - infinity to infinity del L del theta whole square into f x theta dx.

So, by definition, this is the variance of theta hat and this part is the Fisher information statistics, so, variance of theta hat into I theta so this because this expression is = 1, therefore, 1 must be less than equal to variance of theta hat into I theta. So, from this now we will get that variance of theta hat is greater than = 1 by I theta.

(Refer Slide Time: 21:42)

Proof...

- ❖ Thus, we have

$$\Rightarrow \text{var}(\hat{\theta}) \geq \frac{1}{I(\theta)}$$
- ❖ The equality will hold when

$$\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} \sqrt{f(\mathbf{x}; \theta)} = c(\hat{\theta} - \theta) \sqrt{f(\mathbf{x}; \theta)}$$
 implying when $\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} = c(\hat{\theta} - \theta)$
 where c is independent of $\hat{\theta}$ and may be a function of θ .
- ❖ To find c ,

$$E\left(\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta}\right)^2 = c^2 E(\hat{\theta} - \theta)^2$$

$$c^2 = \frac{E\left(\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta}\right)^2}{E(\hat{\theta} - \theta)^2} = \frac{I(\theta)}{\text{var}(\hat{\theta})}$$

So does we have variance of theta hat is greater than = 1 by I theta the equality will hold when 1 of the function that is suppose pulse function del L x theta del theta into root over f x theta that must be multiple of the other functions 3 times theta hat - theta into root over f x theta that is the condition for equality indicates of Cauchy Schwartz inequality. So, since this quantity, actually

we want this to be greater than 0 in the region of support, that is the regularity condition therefore, we can cancel this to quantity.

So what we will get this equality will hold when $\frac{\partial L}{\partial \theta} = c(\theta \hat{\theta} - \theta)$. This is the condition for equality where c is independent of θ as it is a scalar actually. And may be a function of θ c maybe a function of θ but c does not involve x or c does not involve $\hat{\theta}$. Now to find c if we take this square of this and they take the expectation then E of $\left(\frac{\partial L}{\partial \theta}\right)^2$ will be $= c^2 \times E(\theta \hat{\theta} - \theta)^2$.

So, that way c^2 will be equal to this divided by this E of $\left(\frac{\partial L}{\partial \theta}\right)^2$ whole square divided by $E(\theta \hat{\theta} - \theta)^2$ whole square. This is the variance and this is the important statistics. So, that way $I(\theta)$ divided by variance of θ , so, $c^2 = I(\theta)$ divided by variance of θ . So, this scalar c is given by this.

(Refer Slide Time: 23:42)

Alternate form of $I(\theta)$

$$\int_{-\infty}^{\infty} f(x; \theta) dx = 1 \Rightarrow \int_{-\infty}^{\infty} \frac{\partial}{\partial \theta} f(x; \theta) dx = 0$$

$$\therefore \int_{-\infty}^{\infty} \frac{\partial L(x; \theta)}{\partial \theta} f(x; \theta) dx = 0 \quad \text{or} \quad E \frac{\partial L(x; \theta)}{\partial \theta} = 0$$

❖ Taking the partial derivative with respect to θ again,

$$\int_{-\infty}^{\infty} \left\{ \frac{\partial^2 L(x; \theta)}{\partial \theta^2} f(x; \theta) + \frac{\partial L(x; \theta)}{\partial \theta} \frac{\partial}{\partial \theta} f(x; \theta) \right\} dx = 0$$

$$\therefore \int_{-\infty}^{\infty} \left\{ \frac{\partial^2 L(x; \theta)}{\partial \theta^2} f(x; \theta) + \left(\frac{\partial L(x; \theta)}{\partial \theta} \right)^2 f(x; \theta) \right\} dx = 0$$

$$\therefore I(\theta) = E \left(\frac{\partial L(x; \theta)}{\partial \theta} \right)^2 = -E \frac{\partial^2 L(x; \theta)}{\partial \theta^2}$$

Now, we know what is $I(\theta)$ but there is a better expression algebraically for $I(\theta)$ that we can derive now, that is integration of $f dx = 1$, then PDF likelihood function degraded to dx must be $= 1$. This implies that partial derivative of this quantity must be $= 0$. Again, under the regularity conditions, we can take this partial derivative inside therefore $\frac{\partial}{\partial \theta} \int f(x; \theta) dx = 0$. Integrating from $-\infty$ to $+\infty$ must be $= 0$.

Now this we can write as this quantity, now we can write in terms of log likelihood function. So that way $\int_{-\infty}^{+\infty} \frac{\partial L}{\partial \theta} f(x) dx$ that must be $= 0$. This is the expression we get from here, just violating this expression in terms of log likelihood function as we did earlier. So, now, what does this expression implies that expected value of the partial derivative of log likelihood function is always $= 0$ this is also 1 important result.

But we are interested to the second partial derivative therefore, will again take the partial derivative of this expression that way will get, because now it is a product of 2 functions. So, if we take the partial derivative with respect to θ again and of course then that this partial derivatives exist. So, second order partial derivative of L with respect to θ into $f(x)$ that is a part term plus again this term is into partial derivative of this so a known this again we will be writing in terms of log likelihood function that way we will get this square here into $f(x)$ θ dx .

So, from this expression now, if I consider this is nothing, but the expected value of this quantity and if I consider this is the expected value of this quantity. So, that way what I get is that because I know that expected value of this quantity this quantity is the information so, therefore, $I(\theta) = -E$ of this quantity and that must be equal to if I take this quantity and said it will be negative, so that way $-E$ of $\frac{\partial^2 L}{\partial \theta^2}$.

So, $-E$ of $\frac{\partial^2 L}{\partial \theta^2}$ so that way this information statistic $I(\theta)$ can be expressed in terms of the second order partial derivative that is the whereas value of the second order partial derivative with a negative side and compared to taking the square of the partial derivative and then expected value this expected value is easier to compute.

(Refer Slide Time: 27:28)

Therefore,

❖ CRLB of $\text{var}(\hat{\theta})$ is

$$\frac{1}{I(\theta)} = \frac{1}{E \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta^2}}$$

❖ CRLB is reached if

$$\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} = c(\hat{\theta} - \theta)$$

So, therefore, you can write Cramer Rao lower bound for variants of theta = 1 by I theta that we have already established that = $-\frac{1}{E \frac{\partial^2 L}{\partial \theta^2}}$. So, that way we can simplify this expression and also we have a establish is that CRLB is released if this part of that partial derivative del L del theta = c times theta hat - theta. So, this is the Cramer Rao lower bound given by this expression and this CRLB with this, if this likelihood function satisfy this relationship.

(Refer Slide Time: 28:11)

Remarks

- ❖ CRLB specifies the optimal bound that we can hope to achieve in terms of the variance of an unbiased estimator. However, this bound may not always be achieved.
- ❖ If $\hat{\theta}$ satisfies CRLB with equality, then $\hat{\theta}$ is an MVUE and also called an efficient estimator.
- ❖ We have derived the CRLB using continuous RVs. The same can be derived for discrete RVs by replacing the PDF by the PMF and using the same regularity conditions.

Let us make some important remarks CRLB specifies the optimal bound that we can hope to achieve in terms of the variance of an unbiased estimator, it is an optimal bound. However, this bond may not always be achieved that is a now we may not release this bound if theta hat

satisfies CRLB suppose we have an estimator with reset this CRLB then $\hat{\theta}$ is an MVUE because its variance will be minimum we cannot have variance lower than that so, $\hat{\theta}$ in that case is an MVUE and also called an efficient estimator. That we have defined earlier we have derived CRLB we using continuous random variables.

So, because we derived assume in joint PDF, but the same can be derived for discrete random variables by replacing the PDF by PMF. And using this same regularity conditions so that way, CRLB is more general than what we derive here. It is true for both continuous and discrete random variables.

(Refer Slide Time: 29:33)

Remarks...

❖ Suppose $X_1, X_2, \dots, X_N$ are iid. Then

$$I_1(\theta) = E \left(\frac{\partial}{\partial \theta} \ln(f(x; \theta)) \right)^2 = -E \frac{\partial^2}{\partial \theta^2} \ln(f(x; \theta))$$

$$\therefore I(\theta) = -E \left(\frac{\partial^2}{\partial \theta^2} \ln(f(x_1, x_2, \dots, x_N; \theta)) \right)$$

$$= -E \left(\frac{\partial^2}{\partial \theta^2} \ln \left(\prod_{i=1}^N f(x_i; \theta) \right) \right)$$

$$= -E \left(\frac{\partial^2}{\partial \theta^2} \sum_{i=1}^N \ln(f(x_i; \theta)) \right) = N I_1(\theta)$$

Further suppose $X_1, X_2, \dots, X_N$ are iid independent and identically distributed then suppose if I consider only 1 random variable $I_1(\theta) = E$ of by definition, partial derivative of likelihood function whole square and that is also $= -E$ of $\frac{\partial^2}{\partial \theta^2}$ of log likelihood function. So that we have already established this therefore, if there is only 1 variable here $I_1(\theta)$ I can write log this.

Now $I(\theta)$ is a log of satisfies we have to take the $\frac{\partial^2}{\partial \theta^2}$ of log likelihood function of N random variables. So, now this joint PDF we can write in terms of marginal PDF like this like product of N marginal PDF. So, that way this will be $\frac{\partial^2}{\partial \theta^2}$ of log of product of N marginal PDF $f \times I(\theta)$. And this product now we can because it is logarithm of

the product we can write it as this term of logarithm so that way del square del theta Square of summation i going from 1 to n log likelihood functions $\ln f(x_i; \theta)$.

Now this we can take inside this. So that way and all are identical all random variables are identical therefore they will have the same in summation that is $I_1(\theta)$ therefore, it will be and there are N such as variables therefore it will be N times $I_1(\theta)$ so that way if X sides are iid independent and identically distributed, then and the pisser information statistic part is N random variables equal to N times the information statistic per single random variable. So, that is the result for iid this otherwise, if only independent then we have to add all this all in permission we have to it that is the additivity property of information.

(Refer Slide Time: 32:08)

Example:

Let $X_1, X_2, \dots, X_N$ be iid Gaussian with known variance σ^2 and unknown mean μ . Suppose $\hat{\mu} = \frac{1}{N} \sum_{i=1}^N X_i$ which is unbiased. Find CR bound and hence show that $\hat{\mu}$ is an efficient estimator.

Likelihood function

$$f(x_i; \mu) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \left(\frac{x_i - \mu}{\sigma} \right)^2}$$

$$L(x_i; \mu) = -\ln(\sqrt{2\pi\sigma^2}) - \frac{1}{2} \left(\frac{x_i - \mu}{\sigma} \right)^2$$

$$\frac{\partial^2 L(x_i; \mu)}{\partial \mu^2} = -\frac{1}{\sigma^2} \Rightarrow I_1(\theta) = \frac{1}{\sigma^2}$$

$$I(\theta) = N I_1(\theta) = \frac{N}{\sigma^2}$$

$$\therefore \text{CRLB} = \frac{1}{I(\theta)} = \frac{1}{\frac{N}{\sigma^2}} = \frac{\sigma^2}{N}$$

We will give 1 example suppose $X_1, X_2, \dots, X_N$ are iid Gaussian with known variance σ^2 and unknown mean μ . Suppose $\hat{\mu} = \frac{1}{N} \sum_{i=1}^N X_i$, i going from 1 to N divided by N this is the sample mean which is unbiased to have already proved. Find CR bound enhance so, that $\hat{\mu}$ is an efficient estimator. So, f is the likelihood function log likelihood function for single variable it will be like this we are taking the logarithm therefore, it will be - half this because it here, 1 by is there and then - half of $x_1 - \mu$ by σ whole squared now we have to take this second partial derivative.

So, only this term will involve $x - \mu$ therefore, we take the partial derivative and then again derivative will get this = - 1 by σ^2 . Therefore, second order partial derivative $\frac{\partial^2 L}{\partial \mu^2}$

square will be $= -1$ by sigma square because this is this expression only involved mu if you take twice derivative then it will be so, mu will not be there, so, it will be simply 1 divided by sigma squared.

Therefore, $I_1(\theta)$ if I consider only 1 random variable x_1 then its information will be $= 1$ by sigma squared and since exercises are iid therefore, $I(\theta)$ will be N times $I_1(\theta)$ that is equal N by sigma square therefore CRLB Cramer Rao lower bound $= 1$ by $I(\theta)$ that $= 1$ by N by sigma squared, that is sigma square by N . So, this is the Cramer Rao lower bound.

(Refer Slide Time: 34:14)

Example...

Let $X_1, X_2, \dots, X_N$ be iid Gaussian with known variance σ^2 and unknown mean μ . Suppose $\hat{\mu} = \frac{1}{N} \sum_{i=1}^N X_i$ which is unbiased. Find CR bound and hence show that $\hat{\mu}$ is an efficient estimator.

Likelihood function

$$f(\mathbf{x}; \mu) = \prod_{i=1}^N \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(x_i - \mu)^2}{\sigma^2}}$$

$$L(\mathbf{x}; \mu) = -\frac{N}{2} \ln(\sqrt{2\pi\sigma^2}) - \frac{1}{2} \sum_{i=1}^N \left(\frac{x_i - \mu}{\sigma} \right)^2$$

$$\frac{\partial L(\mathbf{x}; \mu)}{\partial \mu} = \sum_{i=1}^N \left(\frac{x_i - \mu}{\sigma} \right)$$

$$= \frac{N}{\sigma} \left(\sum_{i=1}^N \left(\frac{x_i - \mu}{N} \right) \right)$$

$$= \frac{N}{\sigma} (\hat{\mu} - \mu)$$

$\therefore \hat{\mu}$ reaches CRLB

Same example, we are considering so, now, again we consider this likelihood function, this is given by this and if we take the log likelihood function is given by this and now, let us take the first derivative first partial derivative. So, this quantity will be equal to and that is $\frac{\partial L}{\partial \mu}$ will be equal to summation because we will take the partial derivative of this 2 and 2 will get cancel so, that was summation $x_i - \mu$ by sigma i going from 1 to N that is the partial derivative of the log likelihood function.

So, now here sigma is there this sigma I will take out and I will write N here and divided by N here. So, that this expression can be simplified as N by sigma into $\hat{\mu} - \mu$. So, we have seen that $\frac{\partial L}{\partial \mu}$ that is the first partial derivative of log likelihood function is a product of 1 constant term and by sigma into $\hat{\mu} - \mu$. Therefore, according to the equality condition of Cramer Rao theorem $\hat{\mu}$ will receive reaches CRLB. So, because this partial derivative is for

log of this specter mu hat - mu and then scalar therefore mu hat reaches CRLB and hence mu hat is an efficient estimator.

(Refer Slide Time: 35:55)

Summary

- ❖ The uniqueness of MVUE makes it the most desirable
- ❖ $L(\mathbf{x}; \theta) = \ln f(x_1, x_2, \dots, x_n; \theta)$ is called the *log-likelihood function* which characterizes the observed data .
- ❖ $I(\theta) = E\left(\frac{\partial L}{\partial \theta}\right)^2$ is the Fisher information statistic and also related as
- ❖ $I(\theta) = -E\frac{\partial^2 L}{\partial \theta^2}$
- ❖ Cramer Rao theorem- for an unbiased estimator $\hat{\theta}$ under certain regularity conditions,

$$\text{var}(\hat{\theta}) \geq \frac{1}{I(\theta)}$$

Let us summarize the lecture the uniqueness of MVUE makes it the most desirable estimator. So, that uniqueness, we establish $L \times \theta$ that = log of the likelihood function is called the log likelihood function mischaracterizes the observed data $I \theta$ that is the Fisher information statistic that = E of del L del theta squared and also related as $I \theta$ also we expressed theta = – E of del square L del theta square Cramer Rao theorem that also we established for an unbiased estimator theta hat under certain regularity conditions, variance of theta hat is greater than equal to 1 by $I \theta$ is given by this Fisher Information statistics.

(Refer Slide Time: 37:00)

Summary...

❖ CRLB of $\text{var}(\hat{\theta})$ is

$$\frac{1}{I(\theta)} = \frac{1}{E \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta^2}}$$

❖ CRLB is reached if

$$\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} = c(\hat{\theta} - \theta)$$

❖ If $\hat{\theta}$ reaches the CRLB, then $\hat{\theta}$ is an MVUE and an efficient estimator.

Thus CRLB Cramer Rao lower bound of variance of theta hat is 1 by I theta that is the bound and that = - 1 by expected value of del 2 L del theta square CRLB is released if this partial derivative of log likelihood function is a product of c and theta hat - theta where is a constant. If theta hat reaches the CRLB then theta hat is an MVUE so that is very important. So, our MVUE and if theta as reaches Cramer Rao lower bound then it is an MVUE and therefore, it will be an efficient estimator. Thank you.