

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

Lecture No 10
Estimation Theory 3 Cramer Rao Lower Bound II

Hello students, welcome to lecture 10, Cramer Rao Lower Bound two.

(Refer Slide Time: 00:42)

Let us recall

- ❖ The CR theorem gives a bound on the variance of an unbiased estimator in terms of Fisher information statistic.
- ❖ $I(\theta) = E\left(\frac{\partial L}{\partial \theta}\right)^2$ is the Fisher information statistic which can also be expressed as

$$I(\theta) = -E \frac{\partial^2 L}{\partial \theta^2}$$
- ❖ CR theorem- for an unbiased estimator $\hat{\theta}$ under certain regularity conditions,

$$\text{var}(\hat{\theta}) \geq \frac{1}{I(\theta)}$$
- ❖ CRLB for $\text{var}(\hat{\theta})$:

$$\frac{1}{I(\theta)} = \frac{1}{E \left(\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} \right)^2} = - \frac{1}{E \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta^2}}$$
- ❖ CRLB is reached if

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

Let us recall, the Cramer Rao theorem gives a bound on the variance of an unbiased estimator in terms of Fisher information statistics. Now what is a Fisher information statistic? $I(\theta)$ that is expected value of $\frac{\partial L}{\partial \theta}$ squared is the Fisher information statistic, which can also be expressed as $I(\theta) = -E \frac{\partial^2 L}{\partial \theta^2}$.

CR theorem, or Cramer Rao theorem for an unbiased estimator $\hat{\theta}$ under certain regularity conditions, variance of $\hat{\theta}$ is greater than or equal to $\frac{1}{I(\theta)}$. So this is the statement of Cramer Rao theorem. Therefore, Cramer Rao lower bound theorem CRLB for variance of $\hat{\theta}$ is equal to $\frac{1}{I(\theta)}$, that is equal to $\frac{1}{E \left(\frac{\partial L}{\partial \theta} \right)^2}$ and this can be altered to $-\frac{1}{E \frac{\partial^2 L}{\partial \theta^2}}$.

And one important point is CRLB is reached if the partial derivative of log likelihood function that is $\frac{\partial L}{\partial \theta}$ can be expressed as $c(\hat{\theta} - \theta)$, versus θ can be a function of θ .

(Refer Slide Time: 02:16)

This lecture will discuss

- (1) CRLB for estimators of function of a parameter- $g(\theta)$
- (2) CRLB for estimators of vector parameters-

$$\begin{bmatrix} \theta_1 \\ \vdots \\ \theta_k \end{bmatrix}$$

In this lecture we will discuss Cramer Rao lower bound theorem CRLB for estimator of function of a parameter $g(\theta)$ earlier, we discussed the estimator θ but now we will discuss what will be the CRLB for an unbiased estimator for $g(\theta)$. Second thing will be discusses CRLB for estimator of vector parameters.

So, parameters are supposed multiple parameters are there is that θ_1 up to θ_k which are emerged as vector. So what will be there will be CRLB for unbiased estimators for this vector parameter.

(Refer Slide Time: 02:52)

Condition for achieving CRLB

- ❖ Let us revisit the equality condition of the CR theorem:
$$\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} = c(\hat{\theta} - \theta)$$
- ❖ Squaring and taking expectation
$$E\left(\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta}\right)^2 = c^2 E(\hat{\theta} - \theta)^2$$
$$c^2 = \frac{E\left(\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta}\right)^2}{E(\hat{\theta} - \theta)^2} = \frac{I(\theta)}{\text{var}(\hat{\theta})}$$
- ❖ At equality, $\text{var}(\hat{\theta}) = \frac{1}{I(\theta)}$
$$\therefore c^2 = I^2(\theta) \Rightarrow c = I(\theta)$$
- ❖ Thus CRLB is reached if $\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} = I(\theta)(\hat{\theta} - \theta)$

Let us revisit the condition for achieving CRLB, we know that for CRLB we have the equality condition, given by the $\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta}$ is equal to c times $\hat{\theta}$ minus θ .

Squaring and taking expectation we get the sum of $\frac{\partial L}{\partial \theta}$ whole squared is equal to c squared times θ head minus θ whole squared. So we will square is the expectation of both sides.

So that way we get this from this we get, c squared is equal to E of $\frac{\partial L}{\partial \theta}$ Squared, that is $I(\theta)$ divided by up to that θ head minus θ squared. That is called $I(\theta)$. So this is the expression for C squared, which we derived in the last lecture. Now, at the equality, we know that variance of θ head is equal to one by $I(\theta)$. So if I put variance of θ is equal to one by $I(\theta)$.

Then, I will get c squared is equal to $I(\theta)$ from this c square is equal to variance of $I(\theta)$ and will interpret that is equal to one by $I(\theta)$, that the condition, Therefore CRLB is reached if $\frac{\partial L}{\partial \theta}$ equal to $I(\theta)$ into θ head minus θ .

(Refer Slide Time: 04:28)

CRLB for estimators of function of a parameter

- ❖ Suppose $g(\theta)$ is a function of the parameter θ and $\hat{g}(\theta)$ is an unbiased estimator of $g(\theta)$. Then, under the same regularity conditions,

$$\text{var}(\hat{g}(\theta)) \geq \frac{(g'(\theta))^2}{I(\theta)}$$

- ❖ The equality in the CR bound is reached if

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

- ❖ The proof of the above is similar as the proof of the CR theorem discussed earlier.

We will discuss about CRLB for the estimators of function of a parameter. So suppose $g(\theta)$ is a function of the parameter θ and the $\hat{g}(\theta)$ is an unbiased estimator of $g(\theta)$. Then under the same regularity conditions, a variant of the head θ is greater than and equal to the g dash θ whole square divided by $I(\theta)$, that it is the first order derivative with respect to θ .

So this is the Cramer Rao theorem for the function of a parameter and the equality in CR bound for this inequality, the equality will be reached if partial derivative of that is $L(\mathbf{x}; \theta)$, that is partial derivative of the likelihood function is equal to $I(\theta)$ multiplied by the g head

theta minus g theta. So you can factorize del L, del theta to form I theta into g head theta minus g theta.

So this is the condition for reaching the CR bound, the proof of the above is similar to the proof of CR theorem, we have discussed earlier.

(Refer Slide Time: 05:51)

Example:

Let $X_1, X_2, \dots, X_N$ be iid Poisson RVs with unknown parameter λ . Suppose $g(\lambda) = \lambda^2$. Find the CR bound and hence examine if it is achieved.

Likelihood function for X_1

$$p(x_1; \lambda) = \frac{e^{-\lambda} \lambda^{x_1}}{x_1!}$$

Thus, $L(x_1; \lambda) = -\lambda + x_1 \ln \lambda - \ln x_1!$

$$\therefore \frac{\partial^2 L(x_1; \lambda)}{\partial \lambda^2} = -\frac{x_1}{\lambda^2} \Rightarrow I_1(\lambda) = \frac{EX_1}{\lambda^2} = \frac{1}{\lambda}$$

$$\therefore I(\lambda) = NI_1(\lambda) = \frac{N}{\lambda}$$

$$\therefore \text{CRLB} = \frac{(g'(\lambda))^2}{I(\lambda)} = \frac{4\lambda^2}{\frac{N}{\lambda}} = \frac{4\lambda^3}{N}$$

We will consider one example; Let X_1, X_2 , up to X_n be iid poisson random variables with unknown parameter lambda. Suppose the g lambda is equal to lambda square. So find the CR bound for the unbiased estimator of the g lambda. And hence, examine if it is achieved. Now exercise the iid we will find out the likelihood function of x_1 , that is equal to PMF, P_{x_1} is the function of lambda that is equal to e to the power minus lambda into lambda to the power x_1 divided by x_1 factor.

Thus the likelihood function of x_1 will be equal to will take the logarithm, therefore it will be minus lambda plus x_1 in the log of lambda minus log of x_1 and factor in. And if we take the second order partial derivative with respect to lambda will get delta squared is equal to minus x_1 and y lambda square. So, for information, we have to take the expected value, therefore x_1 of lambda will be E of X_1 divided by lambda squared, that is equal to one by lambda.

Since there are n iid random variable, therefore I lambda will be equal to N times I one lambda that is equal to n by lambda. Therefore, CRLB is equal to the g dash lambda squared divided by I lambda that is equal to 4 lambda square. Because if we take the derivative this will be two lambda square, will be 4 lambda squared divided by N lambda, and that is equal

to $2\lambda^3$ by N , and this is a serial way for any unbiased estimator of the λ .
(Refer Slide Time: 07:59)

Example....

Equality condition

✧ Likelihood function

$$p(\mathbf{x}; \lambda) = \prod_{i=1}^N \frac{e^{-\lambda} \lambda^{x_i}}{x_i!}$$

Thus, $L(\mathbf{x}; \lambda) = -\lambda + \ln \lambda \sum_{i=1}^N x_i + \text{terms not involving } \lambda$

$$\therefore \frac{\partial L(\mathbf{x}; \lambda)}{\partial \lambda} = -1 + \frac{\sum_{i=1}^N x_i}{\lambda} = \frac{N}{\lambda} \left(\frac{1}{N} \sum_{i=1}^N x_i - \lambda \right) = I(\lambda)(\hat{\lambda} - \lambda)$$

$\therefore$ CRLB is not reached by an unbiased estimator of λ^2 .

CRLB is achieved by the unbiased estimator of λ given by

$$\hat{\lambda} = \frac{1}{N} \sum_{i=1}^N x_i$$

Now let us examine the equality condition, the likelihood function, we are again relating this P of X vector λ , has a function of λ is equal to product of e to the power minus λ , λ to the power x_i divided by factor $x_i!$, i going from one to N , and this is the joint probability or likelihood function in terms of λ , and taking the logarithm $L(\mathbf{x}; \lambda)$ will be equal to minus λ plus log of summation x_i , i either from one to N plus terms, not involving λ .

Therefore, $\frac{\partial L}{\partial \lambda}$ will be equal to minus one plus derivative of this is minus one plus summation x_i , i going from one to N , and derivative of $\ln N$ is one by λ . So, that this is $\frac{\partial L}{\partial \lambda}$ is equal to minus one plus summation x_i , i going from one to N divided by λ , and this can be simplified as N by λ into summation x_i , i going from one to N minus λ , and this we can write as, $I(\lambda)(\hat{\lambda} - \lambda)$ here it is λ not λ^2 .

Therefore, we conclude that CRLB is not reached by an unbiased estimator of λ^2 . It is reached by an unbiased estimator of λ , which is given by $\hat{\lambda}$ is equal to one by N times summation x_i , i going from one to N .

(Refer Slide Time: 09:48)

CRLB for estimators of vector parameters

- ❖ Suppose $\theta_1, \theta_2, \dots, \theta_k$ are k parameters which are represented as the vector $\mathbf{\theta} = [\theta_1 \dots \theta_k]'$ and characterises the likelihood function

$$f(\mathbf{x}; \mathbf{\theta}) = f(x_1, x_2, \dots, x_n; \theta_1, \theta_2, \dots, \theta_k).$$

- ❖ Then the log-likelihood function is given by

$$L(\mathbf{x}; \mathbf{\theta}) = \ln f(x_1, x_2, \dots, x_n; \theta_1, \theta_2, \dots, \theta_k)$$

- ❖ We can represent the 1st-order partial derivatives of $L(\mathbf{x}; \mathbf{\theta})$ as

$$\frac{\partial}{\partial \mathbf{\theta}} L(\mathbf{x}; \mathbf{\theta}) = \begin{bmatrix} \frac{\partial L(\mathbf{x}; \mathbf{\theta})}{\partial \theta_1} & \frac{\partial L(\mathbf{x}; \mathbf{\theta})}{\partial \theta_2} & \dots & \frac{\partial L(\mathbf{x}; \mathbf{\theta})}{\partial \theta_k} \end{bmatrix}'$$

Now we will be considering CRLB for estimator of vector parameters, suppose θ_1, θ_2 up to θ_k , are k parameters which are represented as a vector, this is a K dimensional vector θ given by θ_1, θ_2 up to θ_k transpose this is a column vector. This vector characterizes the likelihood function and likelihood function is given by $F \times \theta$ that is f of x_1, x_2 up to x_n and parameter started θ_1, θ_2 up to θ_k .

Then the log-likelihood function is given by we will take a logarithm so this is the log-likelihood function. Now because it is a function of several variables or we can get the partial derivative vector. So, we can take the $\partial L / \partial \theta_1, \partial L / \partial \theta_2$ like that so one column we can take this column is the vector representation for the partial derivatives of $\partial L / \partial \theta$.

(Refer Slide Time: 10:56)

Fisher information matrix

- ❖ The Fisher Information matrix is given by

$$\mathbf{I}(\mathbf{\theta}) = E \left(\frac{\partial L(\mathbf{x}; \mathbf{\theta})}{\partial \mathbf{\theta}} \left(\frac{\partial L(\mathbf{x}; \mathbf{\theta})}{\partial \mathbf{\theta}} \right)' \right)$$

where E is performed on each element of the matrix.

- ❖ It can be shown that

$$\begin{aligned} \mathbf{I}(\mathbf{\theta}) &= -E \left(\frac{\partial}{\partial \mathbf{\theta}} \left(\frac{\partial L(\mathbf{x}; \mathbf{\theta})}{\partial \mathbf{\theta}} \right) \right) \\ &= -E \begin{bmatrix} \frac{\partial^2 L(\mathbf{x}; \mathbf{\theta})}{\partial \theta_1^2} & \frac{\partial^2 L(\mathbf{x}; \mathbf{\theta})}{\partial \theta_1 \partial \theta_2} & \dots & \frac{\partial^2 L(\mathbf{x}; \mathbf{\theta})}{\partial \theta_1 \partial \theta_k} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 L(\mathbf{x}; \mathbf{\theta})}{\partial \theta_k \partial \theta_1} & \frac{\partial^2 L(\mathbf{x}; \mathbf{\theta})}{\partial \theta_k \partial \theta_2} & \dots & \frac{\partial^2 L(\mathbf{x}; \mathbf{\theta})}{\partial \theta_k^2} \end{bmatrix} \end{aligned}$$

Fisher information matrix, the Fisher information matrix is given by $I(\theta)$ matrix. That is given by $E\left[\frac{\partial L}{\partial \theta} \frac{\partial L}{\partial \theta}^T\right]$, this is a vector, $\frac{\partial L}{\partial \theta}$ transpose this product will be a matrix and then we take the expectation of the individual element of matrix. So E is performed on its element of the matrix. So this is the Fisher information matrix and it can be also saw that $I(\theta)$ is equal to minus expectation of $\frac{\partial^2 L}{\partial \theta \partial \theta}$ of the $\frac{\partial L}{\partial \theta}$ vector.

So that way this will become a matrix now and that matrix is given where minus expectation of $\frac{\partial^2 L}{\partial \theta_1^2}$, $\frac{\partial^2 L}{\partial \theta_1 \partial \theta_2}$, $\frac{\partial^2 L}{\partial \theta_1 \partial \theta_k}$ like that last row will be $\frac{\partial^2 L}{\partial \theta_k \partial \theta_1}$, $\frac{\partial^2 L}{\partial \theta_k \partial \theta_2}$, up to $\frac{\partial^2 L}{\partial \theta_k^2}$. So this is the Fisher information matrix.

(Refer Slide Time: 12:15)

Unbiased Estimators of Vector Parameters

❖ Suppose $\hat{\theta} = [\hat{\theta}_1, \hat{\theta}_2, \dots, \hat{\theta}_k]^T$ be an unbiased estimator of $\theta = [\theta_1, \theta_2, \dots, \theta_k]^T$.
Then, $E\hat{\theta} = \theta$.

❖ The covariance matrix of $\hat{\theta}$ is defined by

$$C_{\hat{\theta}} = E(\hat{\theta} - \theta)(\hat{\theta} - \theta)^T$$

Note

- (1) The diagonal element $c_{j,j}$ of $C_{\hat{\theta}}$ represents $\text{var}(\hat{\theta}_j)$.
- (2) $C_{\hat{\theta}}$ is a positive semi-definite matrix in the sense that for any real vector $Z \neq 0$, the quadratic form $Z^T C_{\hat{\theta}} Z \geq 0$.
- (3) Eigen values of a positive semi-definite matrix are non-negative. All the leading principal minors of a positive semidefinite matrix are non-negative.

$$\begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix}$$

Now will be considering suppose, unbiased estimators for vector parameters. That is our concern, suppose θ head is equal to θ_1 head, θ_2 head up to θ_k head transpose, this is a vector be unbiased estimator of θ , θ equal to θ_1 , θ_2 , up to θ_k transpose. So, then $E \theta$ head equal to θ that is unbiased estimators and the covariance matrix of θ head is defined by C_{θ} that is equal to E of θ minus θ head in to θ minus θ head transpose this is the definition of covariance matrix.

Now, we examine some interesting properties of C , this covariance matrix, the diagonal element $C_{j,j}$ of C_{θ} represents the variance of θ head j . So, the j th diagonal element will be the variance of θ head. Similarly, C_{θ} is a positive semi definite

matrix that is very important, this covariance matrix is a positive semi definite matrix, how do I define positive semi definite matrix for any real k dimensional vector Z not equal to zero.

The quadratic form this quantity $z^T C \theta$ will be greater than equal to zero, this is the definition of positive semi definiteness and if it is greater than zero, then it will be positive definite. Now, test for positive definiteness is Eigen values of positive semi definite matrix are non negative. So, if we get all Eigen values greater than equal to zero then it will be positive definite matrix and all the leading principle minors of positive semi definite matrix are non negative, this is a test for simple case.

Suppose I have some matrix like this a_{11}, a_{12}, a_{13} suppose $a_{21}, a_{22}, a_{23}, a_{31}, a_{32}, a_{33}$. So this is the matrix. Now we have to consider all leading principal minor that this we have to consider on the left hand side, we have to consider minor parts minor will be a_{11} this will be greater than equal to zero second minor will be determinant of this, this, and this would be also greater than equal to zero. And third, leading principle minority determinant, that's also greater than equal to zero. So this is the test for positive semi definiteness.

(Refer Slide Time: 15:11)

CR theorem for vector case

- ❖ Assume that the likelihood function $f(x_1, x_2, \dots, x_n; \theta)$ satisfies the regularity conditions stated earlier. Then, the covariance matrix C_θ of any unbiased estimator $\hat{\theta}$ satisfies:
$$C_\theta - I^{-1}(\theta) \geq 0$$

where $I(\theta)$ is the Fisher information matrix and the inequality with respect to the Zero matrix implies that $C_\theta - I^{-1}(\theta)$ is a positive semi-definite matrix.
- ❖ The equality will hold when
$$\frac{\partial L(x; \theta)}{\partial \theta} = I(\theta) (\hat{\theta} - \theta)$$

Now, we can state this CR theorem for vector case, assume that the likelihood function, this is the likelihood function satisfies the regularity conditions stated earlier. Then, the covariance matrix C_θ of any unbiased estimator $\hat{\theta}$ that is the covariance matrix satisfy this relationship, covariance matrix minus $I^{-1}(\theta)$, θ is the fisher information matrix that will be greater than equal to zero.

So C of θ head minus I inverse of θ , will be always greater than equal to zero matrix. Where I of θ is the Fisher information matrix and the inequality now this inequality the matrix in inequality, this inequality, with respect to zero matrix implies that this quantity θ head minus I inverse of θ is a positive semi definite matrix. So you have to test whether this matrix C of θ head minus I inverse is a positive semi definite matrix.

This is the condition for Cramer Rao Lower Bound. So this C of θ head by the inverse of θ is a positive semi definite matrix. Now again the equality will be coming in the power vector the $\frac{d}{d\theta} L$ of θ head will be equal to I of θ head minus θ . So this is the condition to test the weather this one bound will be reached.

(Refer Slide Time: 16:52)

Remarks

- ❖ The CR theorem for the individual variances is given by

$$\text{Var}(\hat{\theta}_i) \geq I^{-1}(0)_{i,i}$$

- ❖ The above bound is larger than the one obtained by the CRLB for of $\hat{\theta}_i$ when other parameters are known.
- ❖ Finding $I^{-1}(0)$ is difficult except when X_i s are independent and consequently $I(0)$ is diagonal.

We will make some remarks. Now, what will be the individual suppose we have found this bound in terms of C of θ head, but what will be the real bound for individual element of C of θ head, so that the individual variance will be given by variance θ i head will be greater than equal to I diagonal inverse of θ . That is I diagonal elements of I inverse of θ . This is the variance of θ i.

The bound is larger than the one obtained by CRLB for of θ one head when other parameters are known. For example we can find out individually, the bound on variants of θ i head. Assuming that other parameters are known, in that case we will get a lower bound. So, this bound what do we get through this matrix from these matrix, and that is larger than what we get, individually.

That is larger than the one obtained by CRLB for theta i head when other parameters are known. We will illustrate this concept by finding I inverse theta is difficult except when XiS are independent. In that case I theta is diagonal matrix, so finding out I theta inverse is easier.

(Refer Slide Time: 18:25)

Example CRLB for multiple parameters

Suppose $X_n = a + bn + V_n$, $V_n \sim N(0, \sigma^2)$, a and b are unknown constants. Here $\theta = [a \ b]^T$. To find the CRLB for θ .

The likelihood function is given by

$$f(x_1, x_2, \dots, x_N; a, b) = \frac{1}{(\sqrt{2\pi\sigma^2})^N} e^{-\frac{1}{2\sigma^2} \sum_{i=1}^N (x_i - a - bi)^2}$$

$$L(x; \theta) = -\ln(\sqrt{2\pi\sigma^2})^N - \frac{1}{2\sigma^2} \sum_{i=1}^N (x_i - a - bi)^2$$

$$\therefore \frac{\partial L}{\partial a} = \frac{1}{\sigma^2} \sum_{i=1}^N (x_i - a - bi), \quad \frac{\partial L}{\partial b} = \frac{1}{\sigma^2} \sum_{i=1}^N (x_i - a - bi)i,$$

$$\frac{\partial^2 L}{\partial a^2} = \frac{-N}{\sigma^2}, \quad \frac{\partial^2 L}{\partial a \partial b} = \frac{-N(N+1)}{2\sigma^2} \quad \text{and} \quad \frac{\partial^2 L}{\partial b^2} = \frac{-N(N+1)(2N+1)}{6\sigma^2}$$

We will consider an example of CRLB for multiple parameters. Suppose X_n that is equal to $a + bn + V_n$, this is a constant and it is a trend here bn plus V_n . Where V_n is a normal Gaussian noise and noise sequences are independent, a and b are unknown parameters. Here, parameter vector is equal to a, b transpose. Now we have to find out the CRLB for theta, so in this case because these variances are independent.

Therefore, we can write the join PDF that is the likelihood function that will be given by this expression $\frac{1}{(\sqrt{2\pi\sigma^2})^N} e^{-\frac{1}{2\sigma^2} \sum_{i=1}^N (x_i - a - bi)^2}$ this is the N dimensional Gaussian and we can take the logarithm to get the log likelihood function given by this.

This is this expression for log likelihood function and if I take the partial derivative with respect to a will get this quantity. Similarly, $\frac{\partial L}{\partial b}$ will be given by this, because bi is there. So, we will have I here and similarly, we can take the second order derivative now, suppose $\frac{\partial^2 L}{\partial a^2}$ is given though we can find out $\frac{\partial^2 L}{\partial a^2}$, we can find out $\frac{\partial^2 L}{\partial a \partial b}$, $\frac{\partial^2 L}{\partial b^2}$.

(Refer Slide Time: 20:16)

Example...

$$\therefore \mathbf{I} = \frac{1}{\sigma^2} \begin{bmatrix} N & \frac{N(N+1)}{2} \\ \frac{N(N+1)}{2} & \frac{(N+1)(2N+1)}{6} \end{bmatrix}$$

Taking the inverse we get,

$$\therefore \mathbf{I}^{-1} = \sigma^2 \begin{bmatrix} \frac{2(2N+1)}{N(N-1)} & \frac{-6}{N(N-1)} \\ \frac{-6}{N(N-1)} & \frac{12}{N(N^2-1)} \end{bmatrix}$$

$$\therefore \text{var}(\hat{a}) \geq \frac{2(2N+1)\sigma^2}{N(N-1)} \text{ and } \text{var}(\hat{b}) \geq \frac{12\sigma^2}{N(N^2-1)}$$

❖ You can try to find the bound on $\text{var}(\hat{a})$ when b is a known parameter and compare with the above bound.

So, these are the fellows therefore $\mathbf{I}$ will be equal to one by sigma square times the matrix, first element will be N , and then also second element will be N into $N+1$ divided by 2. Similarly, second row, first element is N into $N+1$ divided by 2 and $N+1$ into $2N+1$ divided by 6, this is the information matrix. Taking the inverse, which is a two by two matrix we can easily take the inverse.

The inverse, $\mathbf{I}^{-1}$ is equal to given by sigma square times, or sentiment is 2 into $2N+1$ divided by N into $N-1$, is minus 6 by N into $N-1$, similarly this element and this element is same. This last element is 12 by N into N squared minus 1. Now, this term is related to the variance of parts parameter, it is variance of a head. Similarly this part is related to the variance of the second parameter that is variants of b head is greater than equal to 2 into $2N+1$ time sigma squared by N into $N-1$.

And variance of b head is greater than equal to 12 sigma squared by N into N squared minus 1. You can try to find the bound on variance of a head when b is a known parameter and compare with the above bound, we have found a bound, when both a and b are unknown, but suppose b is known then try to find out the bound on the variance of a head and compare this bound.

(Refer Slide Time: 22:07)

Practical importance of CRLB

- ❖ CRLB sets a bound on the variance of an unbiased estimator. We cannot have an unbiased estimator with variance lower than the CRLB
- ❖ If the CRLB is achieved by an unbiased estimator $\hat{\theta}$, then $\hat{\theta}$ is the MVUE.
- ❖ The CR theorem tells us the exact condition on $\frac{\partial L}{\partial \theta}$ for achieving the CRLB. If this condition is not satisfied, we never achieve it.
- ❖ The CRLB is a benchmark on the performance of an unbiased estimator. While designing an unbiased estimator, we have to ascertain the closeness of its variance to the CRLB.
- ❖ There is a class of estimators, that asymptotically achieves the CRLB.

Practical importance of CRLB, Cramer Rao Lower Bound sets a bound on the variance of an unbiased estimator. We cannot have an unbiased estimator with a variance lower than the CRLB and if CRLB is achieved by an unbiased estimator $\hat{\theta}$, then $\hat{\theta}$ is the MVUE. So that is an important aspect because MVUE will be the CRLB. This CR theorem came tells us the exact condition on $\frac{\partial L}{\partial \theta}$ for achieving the CRLB.

If this condition is not satisfied, we never achieve it. This CRLB is a benchmark on the performance of an unbiased estimator. While designing an unbiased estimator, we have to ascertain the closeness of its variance to the CRLB. There is a class of estimators that asymptotically achieves the CRLB.

(Refer Slide Time: 23:15)

Summary

- ❖ CR theorem: $\text{var}(\hat{\theta}) \geq \frac{1}{I(\theta)}$
-gives the lower bound for the variance of an unbiased estimator
 - ❖ The CRLB is reached iff $\frac{\partial L(\mathbf{x}/\theta)}{\partial \theta} = I(\theta)(\hat{\theta} - \theta)$
 - ❖ CR theorem for the unbiased estimator of a function
 - $\text{var}(\hat{g}(\theta)) \geq \frac{(g'(\theta))^2}{I(\theta)}$
 - The CRLB is achieved iff $\frac{\partial L(\mathbf{x}/\theta)}{\partial \theta} = I(\theta)(\hat{g}(\theta) - g(\theta))$
 - ❖ CR theorem is extended to unbiased estimator $\hat{\theta} = [\hat{\theta}_1 \hat{\theta}_2 \dots \hat{\theta}_p]'$ of a vector parameter $\theta = [\theta_1 \theta_2 \dots \theta_p]'$ in terms of the Fisher information matrix
- $$I(\theta) = E \frac{\partial}{\partial \theta} L(\mathbf{x}; \theta) \left(\frac{\partial}{\partial \theta} L(\mathbf{x}; \theta) \right)'$$

Let us summarize the lecture; first Cramer Rao theorem variance of theta head is greater than equal to one by I theta. It gives the lower bound for the variance of an unbiased estimator. Theta head is an unbiased estimator of theta. This CRLB is reached iff del L del theta equal to I theta into theta head minus theta. Cramer Rao theorem unbiased estimator of a function is given by variants of g head theta is greater than equal to g dash theta whole squared divided by I theta.

Where g dash theta is the derivative of g theta with respect to I theta. This CRLB is achieved iff the del L del theta is equal to I theta into g head theta that estimator minus g theta. This is the condition for receiving the CRLB by g head theta. CR theorem is extended to an unbiased estimator theta head of a vector parameter theta, which is given by theta 1, theta 2, up to theta k transpose and it is given in terms of Fisher information matrix.

It is given in terms of fisher information matrix I theta, this is a matrix. This is given by del L del theta this is a vector into del L del theta transpose. So this matrix is the Fisher information matrix.

(Refer Slide Time: 24:52)

Summary...

- ❖ Fisher information matrix is also given by

$$I(\theta) = -E \begin{bmatrix} \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta_1^2} & \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta_1 \partial \theta_2} & \dots & \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta_1 \partial \theta_k} \\ \vdots & \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta_2 \partial \theta_1} & \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta_2^2} & \dots & \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta_2 \partial \theta_k} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta_k \partial \theta_1} & \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta_k \partial \theta_2} & \dots & \dots & \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta_k^2} \end{bmatrix}$$
- ❖ CR theorem for multiple parameters- the covariance matrix $C_{\hat{\theta}}$ of any unbiased estimator $\hat{\theta}$ satisfies:

$$C_{\hat{\theta}} - I^{-1}(\theta) \geq 0$$
 implying that $C_{\hat{\theta}} - I^{-1}(\theta)$ is a positive semi-definite matrix.
- ❖ The CRLB is reached

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

Fisher information matrix is also given by this expression, I theta matrix is equal to minus expectation of del 2 L del theta 1 square, del 2 L del theta 1 del theta 2, del 2 L, del theta 1, del theta k and going this way, last row will be del 2 L del theta k del theta 1, del 2 L del theta k del theta 2, up to del 2 L, del theta k square.

Cramer Rao theorem for multiple parameters, the covariance matrix C_{θ} of any unbiased estimator $\hat{\theta}$ satisfies this is covariance matrix, minus inverse of information matrix will always greater than equal to zero matrix. This implying that covariance matrix minus inverse of information matrix, is a positive semi definite matrix. And similarly the CRLB is reached if $\frac{\partial L}{\partial \theta}$, this is a vector equal to $I(\theta)(\hat{\theta} - \theta)$ vector.

(Refer Slide Time: 26:08)

❖ If the MVUE reaches the CRLB, it can be obtained through the factorization condition

$$\frac{\partial L(\mathbf{x}/\theta)}{\partial \theta} = I(\theta)(\hat{\theta} - \theta).$$

However, the CRLB may not be achieved by the MVUE.

❖ We will discuss another approach to find MVUE in the next lecture.

If the MVUE reaches the CRLB, it can be obtained through the factorization condition $\frac{\partial L}{\partial \theta}$ is equal to $I(\theta)(\hat{\theta} - \theta)$. However, the CRLB may not be achieved by the MVUE. We will discuss another approach to find MVUE in the next lecture.

Thank you.