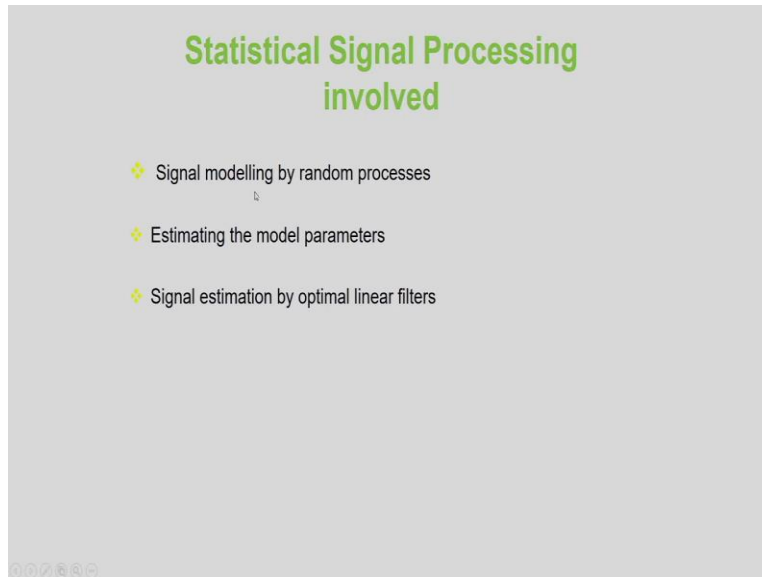


Statistical Signal Processing
Prof. Prabin Kumar Bora
Department of Electronics and Electrical Engineering
Indian Institute of Technology – Guwahati

Lecture - 38
Review 1

(Refer Slide Time: 00:41)



Hello students we are approaching the end of the course in this lecture. We will review some of the important concepts our course involved signal modeling by random processes, estimating the model parameters, signal estimation by optimal linear filters. We will review the first two topics first.

(Refer Slide Time: 01:02)

Total Probability theorem

Let the events $A_1, A_2, \dots, A_n$ form a partition in S so that

$$S = A_1 \cup A_2 \cup \dots \cup A_n \text{ and } A_i \cap A_j = \emptyset \text{ for } i \neq j.$$

Then for any event B ,

$$P(B) = \sum_{i=1}^n P(A_i)P(B / A_i)$$

Bayes rule is given by

$$P(A_k | B) = \frac{P(A_k)P(B / A_k)}{\sum_{i=1}^n P(A_i)P(B / A_i)} \quad k = 1, 2, \dots, n$$

We will start with some basic probability concepts first one is total probability theorem suppose a 1 a 2 up to a n form a partition on the sample space S. So, a 1 a 2 may be up to a n this form a partition on the sample space as that means union of all these events is the sample space S and they do not have any intersection among them. Then any event being suppose being an event the probability of event B we establish that this is equal to summation P of A_i into P of B given A_i , i going from 1 to n.

Now using the total probability theorem we can derive the Bayes rule that is probability of A k given B suppose some event B has occurred what is the probability that a particular event A k occurs so that is the probability of A k given B and this is given why this result probability of A k into probability of B given A k divided by this probability of B is obtained to the total probability theorem. So, that way Bayes theorem is important for us here this probability of all those events with form a partitions these probabilities are known as the apriori probability or prior probability.

And this is the conditional probability then using this apriori probability and conditional probability or the be given A k we can find the Bayes rule. So, Bayes rule will determine the posterior probability, probability of A k given B and this we can determine using this formula.

(Refer Slide Time: 03:09)

➤ A random variable X is characterised by
 CDF $F_X(x) = P(\{s \mid X(s) \leq x\})$,
 PMF $p_X(x) = P(\{s \mid X(s) = x\})$ (discrete case)
 PDF $f_X(x)$ given by $F_X(x) = \int_{-\infty}^x f_X(u) du$ (continuous case)

• The joint RVs X and Y are characterised by joint CDF $F_{X,Y}(x,y)$, joint PMF $p_{X,Y}(x,y)$
 and joint PDF $f_{X,Y}(x,y)$.

➤ X and Y are independent iff

$$F_{X,Y}(x,y) = F_X(x)F_Y(y) \quad \forall x,y \in \mathbb{R}^2$$

Equivalently iff $f_{X,Y}(x,y) = f_X(x)f_Y(y) \quad \forall x,y \in \mathbb{R}^2$

Next we define a random variable you know that random variable is characterized by a CDF F_X of x later on we simplified the notation we simply write this is equal to F of x . This is the probability of the event s such that $X(s)$ is less than equal to small x . Similarly if X is a discrete random variable we define the probability mass function small p_X x is equal to probability of those s for which $X(s)$ is equal to small x , probability of s such that $X(s)$ is equal to X .

And in the case of discrete case we define the PDF. So, PDF small f_X and here the CDF is related to the PDF through this integration. So, PDF at X CDF at X is the integration of $F_X(u) du$ u going from $-\infty$ to x , so this was the definition of PDF. Joint random variable x and y are characterized by the joint CDF that is capital $F_{X,Y}$ at point xy or joint PMF small $p_{X,Y}$ at point xy or joint PDF small $f_{X,Y}$ at point xy .

Later on we simplified these notations so we do not simply $F_{X,Y}$ here similarly this is $p_{X,Y}$ and this one is small $f_{X,Y}$ and the concept of independence is important 2 events 2 random variables X and Y are independent if the joint CDF is product of the marginal CDF. So, $F_{X,Y}$ is equal to F_X into F_Y for all xy belonging to $\mathbb{R}^2$, so this was the definition similarly in terms of joint PDF also we can define it. So, joint PDF f point xy is the product of the marginal PDF. So that way we define the independence.

(Refer Slide Time: 06:11)

Mean, variance and covariance matrix

- The expectation of a function $Y = g(X)$ is given by $Eg(X) = \int_{-\infty}^{\infty} g(x)f_X(x)dx$,
- Correlation of X and Y is given by $E(XY)$, covariance $Cov(X,Y) = E(X - \mu_X)(Y - \mu_Y)$
- For a random vector $\mathbf{X}$, we have

Mean vector $\boldsymbol{\mu}_X = \begin{bmatrix} EX_1 \\ EX_2 \\ \vdots \\ EX_N \end{bmatrix}$

correlation matrix $\mathbf{R}_X = E\mathbf{X}\mathbf{X}^T$ and the covariance matrix $\mathbf{C}_X = E(\mathbf{X} - \boldsymbol{\mu}_X)(\mathbf{X} - \boldsymbol{\mu}_X)^T$

- For the Gaussian random vector $\mathbf{X} \sim N(\boldsymbol{\mu}_X, \mathbf{C}_X)$

$$f_X(\mathbf{x}) = \frac{e^{-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu}_X)^T \mathbf{C}_X^{-1}(\mathbf{x} - \boldsymbol{\mu}_X)}}{(\sqrt{2\pi})^N \sqrt{\det(\mathbf{C}_X)}}$$

- The conditional PDF and the conditional PMF are used to define the conditional expectation $E(Y | X = x)$

So next we saw how we can characterize random variables in terms of its mean variance moments etc. The expectation of a function Y is equal to gX is given by this formula $E gX$ is equal to integration $gx f_X dx$ x going from $-\infty$ to $+\infty$. So, this is the definition of expectation of any function of gX . So, we can use this definition to find out μ is equal to E of X so we will put gX is equal to X here.

E of x square then variance σ^2 is equal to E of $x - \mu$ whole square so we put gX is equal to $X - \mu$. So, that way we can find out various expectations given the PDF or PMF and for two random variables x and y we have the joint expectation $E xy$. So, it is also called correlation and similarly covariance of xy is E of $x - \mu_x$ into $y - \mu_y$ this is the function covariance. And if we have two or more random variables we can represent the random variable by a vector.

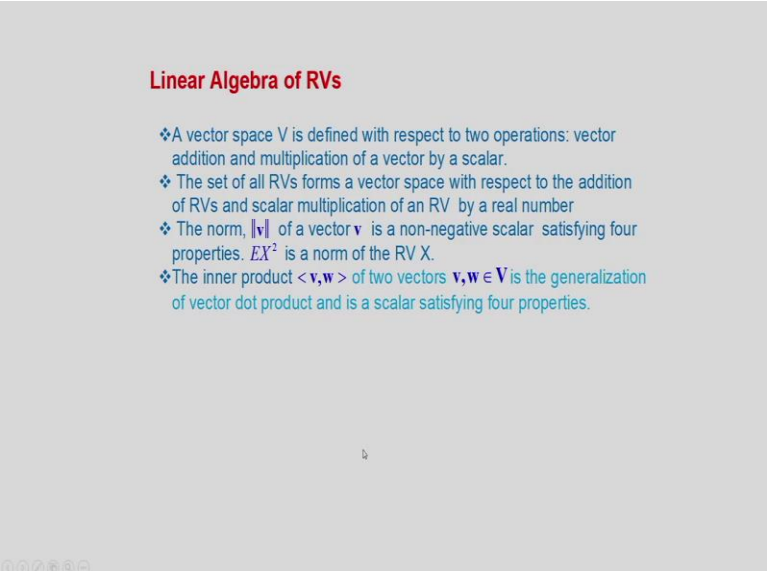
For example this x may be represented as a vector that is there are N random variables like this suppose then this can be represented as a vector. And now given the random vector we can find out the mean vector $\boldsymbol{\mu}_x$, $\boldsymbol{\mu}$ is a vector its components are E of X_1 E of X_2 up to E of X_N . Similarly we can define correlation metrics to measure the correlation among the random variables we define the correlation matrix it is given by E of X given X transpose that means E of this one is the defined as E of that it X is this $X_1 X_2$ up to X_N .

Then transpose will be a row vector $X_1 X_2 \dots X_N$ so this will give us a matrix and then we will take the expectation of individual elements. Then covariance matrix similarly is defined C_x is equal to $E[(X - \mu_x)(X - \mu_x)^T]$. And one important distribution is the Gaussian distribution for multiple random variables we have the Gaussian random vector and this is denoted by normal me did the one by the parameter μ_x vector and the covariance C_x matrix.

And its distribution is given by f_x is equal to $e^{-\frac{1}{2}(x - \mu_x)^T C_x^{-1} (x - \mu_x)}$ divided by $\sqrt{(2\pi)^n \det(C_x)}$, so this is the expression for the multivariate Gaussian or vector Gaussian random vector. And we also defined the conditional PDF for example conditional PDF we define f_x of x given y is equal to f_{xy} divided by f_y provided f_y is not equal to 0.

And once we define the conditional PDF or conditional PMF similarly we can define we can find out the conditional expectation $E[Y \text{ given } X]$ is equal to $\int y f_{y|x}(y|x) dy$ integration from $-\infty$ to ∞ . So, this is the definition of conditional expectation.

(Refer Slide Time: 11:10)



Linear Algebra of RVs

- ❖ A vector space V is defined with respect to two operations: vector addition and multiplication of a vector by a scalar.
- ❖ The set of all RVs forms a vector space with respect to the addition of RVs and scalar multiplication of an RV by a real number
- ❖ The norm, $\|v\|$ of a vector v is a non-negative scalar satisfying four properties. $E[X^2]$ is a norm of the RV X .
- ❖ The inner product $\langle v, w \rangle$ of two vectors $v, w \in V$ is the generalization of vector dot product and is a scalar satisfying four properties.

Then we discussed about linear algebra random variables. We discussed that a vector space V is defined with respect to two operations vector addition and multiplication of a vector by a scalar. The set of all random variables form a vector space with respect to the addition of random

variables and scalar multiplication when random variable of a random variable by a real number this scalar in case of real vector space either is a real number. So, therefore a vector space is defined with respect to two operations here one is the summation of two random variable that is addition and the other one is multiplication of a random variable by a scalar.

Then we define a normal way vector V this is the notation is a non-negative scalars satisfying four properties. So, what are those properties that is norm of X is always greater than equal to 0 number two norm of X is equal to 0 if X is equal to 0, number two if and only if, number three suppose r is a scalar then norm of rx is equal to absolute value of r into norm of x . So, this is for r belonging to real line.

Similarly 4 is norm of $x + y$ suppose x and y are two vectors then norm of $x + y$ is less than equal to norm of x + norm of y . So, norm of $X + y$ is less than equal to norm of x + norm of y this is known as a triangular inequality. So, we also saw that E of X square is the norm of the random variable X we saw that your X square satisfies these four properties so it is a norm. Similarly we define the inner product of two vectors this is the symbol inner product of v w .

So, this also generalization of dot product of ordinary vectors and it is a scalar satisfying 4 properties here also we saw 4 properties are there and those are like; so number one is that is xy inner product of XY is equal to inner product of Y into inner product of Y X this is for all X Y . Then number 2 inner product of X , X is equal to norm of X square, number 3 is if you consider r X , X that will be equal to r times inner product of XX .

Then that number 4 is if I considered a suppose inner product of X , $Y + Z$ that will be equal to inner product of XY + inner product of X Z , so that way we can define the inner product. So, in the case of random variables E of XY is an inner product operation E of XY is an inner product operation that we have solved.

(Refer Slide Time: 15:48)

Linear algebra of RVs...

- ❖ Cauchy Schwarz inequality is given by $|\langle \mathbf{v}, \mathbf{w} \rangle| \leq \|\mathbf{w}\| \|\mathbf{v}\|$
- ❖ The ratio $\rho = \frac{\text{cov}(X, Y)}{\sigma_X \sigma_Y}$ is called the correlation coefficient. The CS inequality was used to prove $|\rho| \leq 1$
- ❖ For uncorrelated RVs X and Y , $\rho = 0$. If X and Y are uncorrelated Gaussian, they are independent also.
- ❖ Two random variables X and Y are called orthogonal if $\langle X, Y \rangle = EXY = 0$
- ❖ Zero-mean uncorrelated RVs are orthogonal and they play important role in the modelling of random signals.

Then we established one important property that is Cauchy-Schwarz inequality the tip which is given by no magnitude of inner product of $\mathbf{v}$ and $\mathbf{w}$ is less than equal to norm of $\mathbf{w}$ into norm of $\mathbf{v}$ this is the Cauchy Schwarz inequality and because of this E of XY suppose if we take E of XY magnitude is there is less than square root of EX square into square root of EY square, so norm of X is EX square root of E of X square norm of Y is square root of EY square.

So, that way magnitude of E of XY ; is less than equal to square root of EX square into square root of EY square that is known as the Cauchy Schwarz inequality in the case of vector space.

(Refer Slide Time: 16:52)

Linear Algebra of RVs

- ❖ A vector space V is defined with respect to two operations: vector addition and multiplication of a vector by a scalar.
- ❖ The set of all RVs forms a vector space with respect to the addition of RVs and scalar multiplication of an RV by a real number
- ❖ The norm, $\|\mathbf{v}\|$, of a vector $\mathbf{v}$ is a non-negative scalar satisfying four properties. $\sqrt{EX^2}$ is a norm of the RV X .
- ❖ The inner product $\langle \mathbf{v}, \mathbf{w} \rangle$ of two vectors $\mathbf{v}, \mathbf{w} \in V$ is the generalization of vector dot product and is a scalar satisfying four properties.

1. $\|\mathbf{x}\| \geq 0$
2. $\|\mathbf{x}\| = 0$, if $\mathbf{x} = 0$
3. $\|\gamma \mathbf{x}\| = |\gamma| \|\mathbf{x}\|$, $\gamma \in \mathbb{R}$
4. $\|\mathbf{x} + \mathbf{y}\| \leq \|\mathbf{x}\| + \|\mathbf{y}\|$
triangle inequality

Square root of E of X square is a norm of the random variable X . So, once we define the norm we can show that this quantity square root of E of X square is a norm which satisfy these four properties. Similarly we define the inner product of two vectors. So, it also satisfy 4 properties number one is inner product of XY is equal to inner product of YX . Number 2 is inner product of X , X is equal to it is a norm square, number 3 inner product of any scalar multiplier r X , Y that is equal to r times norm of r into inner product of XY .

Then number 4 is inner product of $X + Y$, W is equal to inner product of XW plus inner product of YW . So, that way these 4 properties if it is satisfied by a scalar of portion like this then it is an inner product operation. Then we establish the Cauchy Schwarz inequality which is given by mod of inner product of v and w is less than equal to norm of w into norm of v . So, this we can directly apply to establish this relationship that is magnitude of E XY is less than equal to magnitude of E of X square root of E of X square into square root of E of Y square.

So this result directly follows from Cauchy Schwarz inequality and that also resulted in that covariance of magnitude of covariance of XY is less than square root of $\text{Sigma } X^2$ into $\text{Sigma } Y^2$ so that way we get this ratio ρ is equal to covariance of XY divided by $\text{Sigma } X$ into $\text{Sigma } Y$ and it is called the correlation coefficient and we can use the Cauchy Schwarz inequality to prove that mod of ρ that is the correlation coefficient is less than equal to 1.

For uncorrelated random variables X and Y , ρ is equal to 0 so X and Y are uncorrelated, uncorrelated if ρ is equal to 0, so and this will imply that E of XY is equal to $E X$ into $E Y$. If X and Y are uncorrelated Gaussian they are independent also usually uncorrelatedness does not imply independence but if X and Y are uncorrelated Gaussian they are independent also.

Now using the concept of inner product we can define orthogonal random variables two random variables X and Y are called orthogonal if inner product of XY that is equal to E of XY is equal to 0 so this is the definition of orthogonal random variable 0 mean uncorrelated random variables are orthogonal because in that case uncorrelated random variable E of XY is equal to $E X$ into $E Y$ and 0 mean if any of the mean is 0 then this will become 0 therefore it will become orthogonal.

So, zero mean uncorrelated random variables are orthogonal and they play important role in modeling of random signals.

(Refer Slide Time: 21:28)

Random processes

- ❖ We defined a discrete-time random process $\{X(n)\}$ as an indexed family of random variables, where the index set $\Gamma \subseteq \mathbb{Z}$
- ❖ Stationarity plays an important role in analysing a random process. For a strict-sense stationary process $\{X(n)\}$,

$$F_{X(n_1), X(n_2), \dots, X(n_k)}(x_1, x_2, \dots, x_k) = F_{X(n_1+h), X(n_2+h), \dots, X(n_k+h)}(x_1, x_2, \dots, x_k),$$

$\forall k \in \mathbb{N} \text{ and } \forall h, n_1, \dots, n_k \in \Gamma$

- ❖ An RP $\{X(n)\}$ is a wide-sense stationary if $EX(n) = \text{constant}$ and $R_x(m) = EX(n)X(n+m)$ is a function of lag m only

Next we discuss the concept of random process we defined a discrete-time random process X_n and then index family of random variables where the index set τ is a subset of set of integer and characterizing a random process by probability is difficult very complex therefore we define some easier concept that is stationarity and wide-sense stationarity. For a strict sense stationary random process, for a strict sense stationary process X_n the joint distribution at instant n_1 and n_2 up to n_k is same as the joint distribution at instant $n_1 + h$ $n_2 + h$ up to $n_k + h$.

So joint distribution is invariant under the translations of the time axis so this was the definition of the a strict sense stationarity. And then we consider stationarity in a very narrower sense a random process x_n is wide sense stationary if you have X_n is constant and autocorrelation function, now autocorrelation function is defined as $E[X(n)X(n+M)]$ joint expectation of X_n and X_{n+m} is the function of lag m only so it will not depend on this parameter n but it will depend only the difference parameter like parameter m only that way we define the WSS process.

(Refer Slide Time: 23:32)

Random processes

- ❖ A discrete-time WSS random process $\{X(n)\}$ is characterised in terms of the autocorrelation function $R_X(m)$
- ❖ $R_X(m)$ is an even function of m with a maximum at $m=0$.
- ❖ $R_X(m)$ and $S_X(w)$ form a DTFT pair

$$S_X(w) = \sum_{m=-\infty}^{\infty} R_X(m) e^{-j\omega m} \quad -\pi \leq w \leq \pi$$

$$R_X(m) = \frac{1}{2\pi} \int_{-\pi}^{\pi} S_X(w) e^{j\omega m} dw$$

- ❖ The generalized PSD in the z-domain is defined as

$$S_X(z) = \sum_{m=-\infty}^{\infty} R_X(m) z^{-m}$$

A discrete time WSS random process X_n is characterized in terms of the autocorrelation function $R_X(m)$ and we observe that $R_X(m)$ is an even function of m with a maximum at $m=0$. Then we define the power spectral density that in the case of WSS random process that is power spectral density $S_X(\omega)$ this is the average power per frequency this is defined by that is $S_X(\omega)$ is equal to summation $R_X(m)$ into $e^{-j\omega m}$, m going from minus infinity to plus infinity.

So this is the discrete time Fourier transform of the autocorrelation sequence. So, $S_X(\omega)$ is equal to summation $R_X(m)$ into $e^{-j\omega m}$, m going from minus infinity to plus infinity. So, this is this discrete time Fourier transform of the autocorrelation sequence and here ω is defined uniquely from $-\pi$ to π because this function is a periodic function just like any other DTFT function therefore it is uniquely defined for ω lying between $-\pi$ and π .

And autocorrelation is related with power spectral density by this relationship $R_X(m)$ is equal to $\frac{1}{2\pi}$ integration $S_X(\omega) e^{j\omega m} d\omega$, ω going from $-\pi$ to π . So, this way we define power spectral density and how Auto correlation function is related to power spectral density and for our analysis generalized power spectral density is important this is the Z transform of the autocorrelation sequence and given by this formula $S_X(z)$ that is equal to summation $R_X(m) z^{-m}$, m going from minus infinity to infinity.

(Refer Slide Time: 25:54)

Response of a linear system to WSS inputs

❖ For an LTI system with impulse response $h(n)$ and input $x(n)$, the output

$$y(n) = \sum_{k=-\infty}^{\infty} h(k)x(n-k) = x(n) * h(n)$$

❖ In the transform domain.

$$Y(\omega) = H(\omega) X(\omega)$$

$$H(\omega) = \text{DTFT}(h(n)) = \text{Frequency response of the system}$$

❖ For an LTI system with impulse response $h(n)$ and WSS input $X(n)$,

- $X(n)$ and output $Y(n)$ are jointly WSS
- $\mu_y = \mu_x H(0)$
- $R_y(m) = h(m) * h(-m) * R_x(m)$
- $S_y(\omega) = |H(\omega)|^2 S_x(\omega)$
- $S_y(z) = H(z)H(z^{-1})S_x(z)$

Next we discussed about the response of a linear system to WSS inputs. So, we know that for a linear time-invariant system the input-output relationship is this y_n is equal to summation h_k into x of $n - k$ k going from minus infinity to infinity this is the convolution relationship and we denote it by x_n start h_n and in the transform domain we take the discrete-time Fourier transform of both then Y Ω is equal to H Ω into X Ω where each Ω is the DTFT of h_n that is impulse response sequence and it is called the frequency response of the system.

Now in the case of LTI system with WSS input we know deep following that input X_n and the output Y_n are jointly WSS that we established and expected value of the output μ_y that is equal to μ_x times H of 0 what is H 0? H 0 is the frequency response at 0 frequency that is equal to summation of h_n , n going from minus infinity to plus infinity n is equal to h of 0. Then we establish that R_y m that is a top correlation of the output is h_m is convolved with $h_m - m$ then convolved with R_x of m .

So, R_y of m is h_m convolved $h_m - m$ convolved with R_x of m . And in the frequency domain power spectral density S_y Ω is equal to mod of H Ω square into S_x Ω . So, these we get from here Fourier transform of this will be a S_y Ω this one will be H Ω this is a star Ω because it is a $h_m - m$ together a it will be norm of H Ω or magnitude of H

σ^2 into power spectral density of X that is $S_X(\omega)$ so that way as $S_Y(\omega)$ is equal to $|H(\omega)|^2 S_X(\omega)$.

And in this z transform domain we have $S_Y(z)$ is equal to $H(z)$ into $H(z)^{-1}$ into $S_X(z)$ so this is the relationship in the z transform domain.

(Refer Slide Time: 29:14)

White noise and spectral factorization

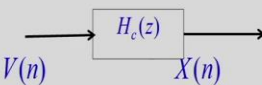
- A white noise process $V(n)$ is characterized by

$$S_V(\omega) = \sigma^2 \quad -\pi \leq \omega \leq \pi$$
 and

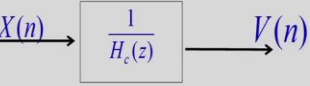
$$R_V(m) = \sigma^2 \delta(m)$$
- $V(n)$ is always zero-mean and samples of $V(n)$ are uncorrelated
- If $V(n)$ is passed through an LTI system, we get non-white WSS output
- The generalized PSD $S_X(z)$ of a regular WSS process $X(n)$ can be factorized as

$$S_X(z) = \sigma^2 H_c(z) H_c(z^{-1})$$

❖ Innovation



❖ Whitening



Then we defined white noise, white noise process $V(n)$ is characterized by constant power spectral density $S_V(\omega)$ is equal to σ^2 ω lying between $-\pi$ and π . So, power spectral density is constant it will imply that autocorrelation function will be $\sigma^2 \delta(m)$ where $\delta(m)$ is the impulse $\delta(m)$ is equal to 1, 1 for m is equal to 0 is equal to 0 otherwise. And what noise is already remain and samples of $V(n)$ are uncorrelated so this implies that samples of $V(n)$ are uncorrelated.

Now one important result if $V(n)$ is passed through an LTI system we get a non-white WSS, non-white, white sense stationary process. So, taking clue from this idea we establish the spectral factorization theorem that is a $S_X(z)$, the power spectral density of any random process or WSS random process under certain condition known as the regularity condition for a regular random process the generalized PSD $S_X(z)$ from this result we establish that the generalized PSD $S_X(z)$ that of a regular WSS process $X(n)$ can be factorized as $S_X(z)$ is equal to σ^2 into $H_c(z)$ into $H_c(z)^{-1}$ where $H_c(z)$ is a minimum phase transfer function.

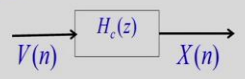
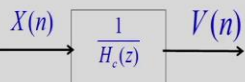
And similarly it is $H_c(z)$ will be a maximum phase and because L transfer function exceeded the minimum phase causal transformation and σ_v^2 is a constant which can be interpreted as the variance of a white noise. And from this result we have the innovation representation of the signal that is $X(n)$ is the output of a linear time-invariant minimum phase system with transfer function $H_c(z)$ which input $V(n)$ so this is the result.

And similarly since this $H_c(z)$ is invertible minimum phase we can pass $X(n)$ through one by a filter up transfer function $1/H_c(z)$ and we will get back $V(n)$ this is known as whitening a widening a WSS process.

(Refer Slide Time: 32:14)

White noise and spectral factorization theorem

- ❖ The generalized PSD $S_X(z)$ of a regular WSS process $X(n)$ can be factorized as

$$S_X(z) = \sigma_v^2 H_c(z) H_c(z^{-1})$$
- ❖ Innovation

- ❖ Whitening

- ❖ Spectral factorization can be used to generate linear models for WSS processes.

Now this is the result important result now spectral factorization can be used to generate linear models for WSS process. So, we can pass $V(n)$ through a filter to get $X(n)$ this is the WSS process therefore WSS process can be expressed in terms of this filter and the input $V(n)$ that way we get the linear models for WSS.

(Refer Slide Time: 32:50)

Linear models...

- ❖ A general ARMA(p,q) model is mathematically described by linear constant-coefficient difference equations.

$$X(n) = \sum_{i=1}^p a_i X(n-i) + \sum_{j=1}^q b_j V(n-j)$$

- ❖ An MA(q) model is an all-zero model given by

$$X(n) = \sum_{i=0}^q b_i V(n-i)$$

- ❖ The ACF $R_X(m)$ is related with the model parameters by

$$R_X(m) = \sum_{j=0}^{q-m} b_j b_{j+m} \sigma_v^2 \quad 0 \leq m \leq q$$

We discussed about linear models is general ARMA pq model is mathematically described by a linear constant-coefficient difference equation that is X_n is equal to summation a_i into X of $n - i$ i going from 1 to p + summation B_j into V of $n - j$ j going from 1 to q so this is the model ARMA p q this is the p term this is the q term, this is known as the auto regression and this is this part is known as the moving average.

And for particular case and MA q model is in all the other model given by X_n is equal to summation b_i into V of $n - i$ only V components are there white noise components are there so it is a linear combination of white noise, it is a linear combination of white noise samples. Now this is the MA q model and its autocorrelation function $R_X(m)$ is related to the model parameters by this formula this we can establish easily that R_X of m is equal to b_j into b_{j+m} σ_v^2 for j going from 0 to $q - m$.

And here m lies between 0 and q and R_X of $-m$ we can find out to be R_X of $m - m$ is also equal to R_X of m therefore here we can write instead of $m \bmod q$. So, this is the autocorrelation function of a moving average process.

(Refer Slide Time: 34:48)

Linear Models

❖ The all-pole AR(p) model is given by

$$X(n) = \sum_{i=1}^p a_i X(n-i) + V(n)$$

❖ The ACF and the PSD are related to the AR parameters by

$$R_X(m) = \sum_{i=1}^p a_i R_X(m-i) + \sigma_v^2 \delta(m), \forall m \in \mathbb{Z}$$

$$S_X(\omega) = \frac{\sigma_v^2}{|H(\omega)|^2} = \frac{\sigma_v^2}{|1 - \sum_{i=1}^p a_i e^{-j\omega i}|^2}$$

❖ ARMA(p,q) process is given by

$$X(n) = \sum_{i=1}^p a_i X(n-i) + \sum_{j=1}^q b_j V(n-j)$$

Similarly we define the all-pole model AR p model this is given by X_n is equal to summation a_i into X of $n - i$ + V_n i doing from 1 to p , so this is the AR p model. The ACF and the PSD are related to the AR parameters by this relationship autocorrelation function is satisfied these relationship R_X of m is equal to summation a_i into R_X of $m - i$ i going from 1 to p + $\sigma_v^2 \delta(m)$ where m is their integer m belong to $\mathbb{Z}$.

So this is the autocorrelation function related with the model parameters. Similarly we can find out the power spectral density also this is $S_{XX}(\omega)$ this is given by σ_v^2 square by mod of $H(\omega)$ Square and this is this we can write σ_v^2 square divided by mod of $1 - \sum_{i=1}^p a_i e^{-j\omega i}$ are going from 1 to p whole square. So, this is the power spectral density of an AR p model.

So, if X_n is a AR p power spectral density is given by this and for an ARMA pq process this is the model there is a p parameter there is a q parameter so that way X_n is equal to summation $a_i x$ $n - i$ i going from 1 to p plus summation $b_j V$ $n - j$ j going from 1 to q so this is the ARMA pq model.

(Refer Slide Time: 36:35)

Linear Models...

❖ The ACF and the PSD are given by

$$R_X(m) = \sum_{i=1}^p a_i R_X(m-i) + \sum_{i=1}^q b_i EV(n+m-i)X(n)$$

$$\text{For } m \geq q+1, \quad R_X(m) = \sum_{i=1}^p a_i R_X(m-i), \quad (\text{Yule Walker equations})$$

$$\text{❖ PSD } S_X(\omega) = \frac{\left| \sum_{i=0}^q b_i e^{-j\omega i} \right|^2}{\left| 1 - \sum_{i=1}^p a_i e^{-j\omega i} \right|^2}$$

And the autocorrelation function and the power spectral density of the ARMA pq model is given by this relationship $R_X(m)$ is equal to summation a_i into $R_X(m-i)$ i going from 1 to p plus the contribution from the moving average part this is given by this summation b_i into $EV(n+m-i)$ into $X(n)$ i going from 1 to q and we observe that for m greater than equal to $q+1$ this $R_X(m)$ will be simply given by this part summation a_i into $R_X(m-i)$ i going from 1 to p this is the Yule walker equation for ARMA p, q model.

And power spectral density since it has two part will be given by $S_X(\omega)$ into numerator is the moving average part summation b_i into $e^{-j\omega i}$ i going from 0 to q whole square summation b_i into $e^{-j\omega i}$ i going from 0 to q mod whole square divided by so PSD $S_X(\omega)$ is given by mod of the numerator polynomial square that is summation b_i into a to the power $j\omega i$ i going from 0 to Q and then mod square.

So this is the numerator part similarly denominator part is contribution from the moving average part and this is given by the norm of $1 - \sum_{i=1}^p a_i e^{-j\omega i}$ whole square so this is the power spectral density for a ARMA pq model. So, that we saw the modeling of random process in terms of linear model needs to discuss parameter estimation.

(Refer Slide Time: 38:54)

Parameter estimation

❖ The observed random data $X_1, X_2, \dots, X_N$ are characterized by a joint

$$\text{PDF } f(x_1, x_2, \dots, x_N; \theta) = f(\mathbf{x}; \theta)$$

❖ An estimator $\hat{\theta}(\mathbf{X}) = \hat{\theta}(X_1, X_2, \dots, X_N)$ is a function by which we guess about the value of the unknown parameter θ .

❖ A particular value $\hat{\theta}(\mathbf{x}) = \hat{\theta}(x_1, x_2, \dots, x_N)$ is called the estimate of the parameter θ .

❖ An estimator $\hat{\theta}$ of θ is unbiased if and only if $E\hat{\theta} = \theta$. The quantity

$$b(\hat{\theta}) = E\hat{\theta} - \theta \text{ is called the bias of } \hat{\theta}.$$

❖ $\text{var}(\hat{\theta}) = E(\hat{\theta} - E\hat{\theta})^2$ is the variance of the estimator

❖ A minimum variance of the unbiased estimator (MVUE) has the lowest variance in the class of unbiased estimators

The parameter estimation problem is like this given the observed random data $X_1, X_2, \dots, X_N$ and the probability model that is $f(x_1, x_2, \dots, x_N)$ as a function of θ this is the probability model and random data is modeled by this probability. In parameter estimation the observed random data $x_1, x_2, \dots, x_N$ are characterized by a joint PDF $F(x_1, x_2, \dots, x_N)$ as a function of θ that is $f(\mathbf{x}; \theta)$.

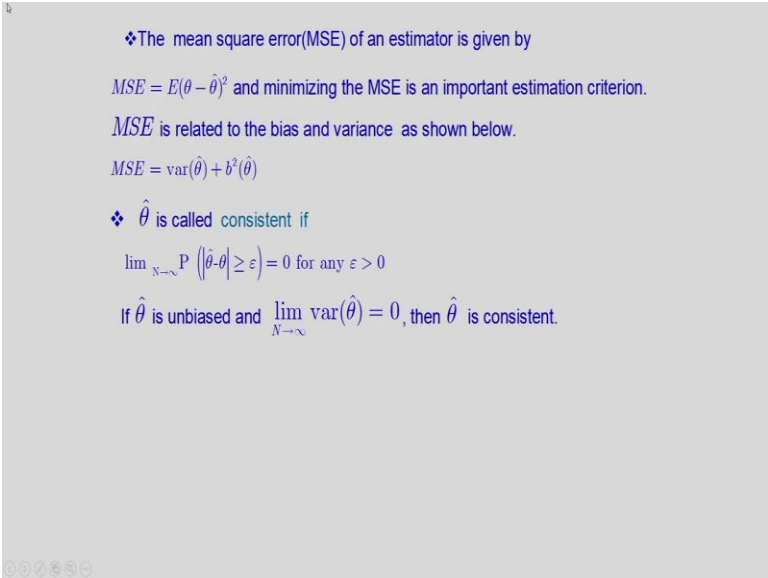
And estimate our $\hat{\theta}$ that is θ it is a function $\hat{\theta}(X_1, X_2, \dots, X_N)$ is a function by which we guess about the value of the unknown parameter θ . So, using this functional relationship we will find out some acceptable value for the unknown parameter θ this we denote it by $\hat{\theta}$. A particular value of the estimator that is $\hat{\theta}(x_1, x_2, \dots, x_N)$ is equal to $\hat{\theta}(x_1, x_2, \dots, x_N)$ is called an estimate of the parameter θ .

So $\hat{\theta}$ that is the estimation rule and here it is an estimator and if we consider particular values of the random vectors that is $x_1, x_2, \dots, x_N$ then we have an estimate of the parameter. So, $\hat{\theta}$ estimate θ we got is the estimator. And estimator $\hat{\theta}$ of θ is unbiased if and only if $E\hat{\theta} = \theta$ that unbiased estimator we define. And if it is a biased estimator we have the bias term $E\hat{\theta} - \theta$ is called bias of the $\hat{\theta}$.

We also define the variance of theta hat it is given by $E(\theta \text{ hat} - E \text{ of } \theta \text{ hat})^2$ that is expected value of theta hat - E of theta hat whole squared this is the variance of theta hat. So, one very desirable properties minimum variance unbiased estimator that is the estimator is unbiased and the variance is lowest among the class of unbiased estimator. So, a minimum variance unbiased estimator MVUE has the lowest variance in the class of unbiased estimator.

So, we define one important property which is most desirable that is minimum variance unbiased estimator MVUE. So, if the estimator is unbiased and it has the variance lowest among the variances in the class of unbiased estimator then it is known as the minimum variance unbiased estimator.

(Refer Slide Time: 42:24)



❖ The mean square error (MSE) of an estimator is given by

$$MSE = E(\theta - \hat{\theta})^2 \text{ and minimizing the MSE is an important estimation criterion.}$$

MSE is related to the bias and variance as shown below.

$$MSE = \text{var}(\hat{\theta}) + b^2(\hat{\theta})$$

❖ $\hat{\theta}$ 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.

We also define the mean square error MSE of an estimator given by $E(\theta - \theta \text{ hat})^2$ that is the MSE and it is related to variance and bias by this relationship MSE is equal to variance of theta hat + biased theta square and some asymptotic property we define theta hat is called consistent if limit probability that theta hat minus theta is greater than any Epsilon this probability goes down to 0 as N tends to infinity that means the probability that theta hat will be deviated from theta by any arbitrary amount as N tends to infinity is equal to 0.

This probability will converge to 0 that way then we will say that theta hat is a consistent estimator. In that case that means theta hat will be very close to the original parameter. Now

particularly if $\hat{\theta}$ is unbiased this definition of consistency can be simplified if $\hat{\theta}$ is unbiased then and limit of variance $\hat{\theta}$ as N tends to infinity is equal to 0 then $\hat{\theta}$ is consistent. So, in the case of unbiased estimator this is the test for consistency variance of $\hat{\theta}$ is equal to 0 as N tends to infinity. Limit of variance of $\hat{\theta}$ is equal to 0 as N tends to infinity.

(Refer Slide Time: 44:14)

MVUE through CRLB

- ❖ The uniqueness of MVUE makes it the most desirable estimator
- ❖ $f(\mathbf{x}; \theta) = f(x_1, x_2, \dots, x_N; \theta)$ is the **likelihood function**.
- ❖ $L(\mathbf{x}; \theta) = \ln f(x_1, x_2, \dots, x_N; \theta)$ 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
- ❖ $I(\theta) = -E\left(\frac{\partial^2 L}{\partial \theta^2}\right)$
- ❖ Cramer Rao theorem- for an unbiased estimator $\hat{\theta}$ under certain regularity conditions,

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

And we thought we established that a very desirable property for estimation is MUE minimum variance unbiased estimator property and we search for this MVUE we first test was true Cramer Rao lower bound. The uniqueness of the MVUE the uniqueness of MVUE makes it the most desirable estimator minimum variance unbiased estimator is the most desirable estimator because it is unique. Now to find the MVUE through CRLB we need the likelihood function that is f of x θ given by this, this is design density as a function of θ and this is called the likelihood function.

And then we defined a log likelihood function log of likelihood function as a function of θ this is called a log likelihood function likelihood function and log likelihood function. they characterize the observed data. Then we define one parameter known as the Fisher information statistic $I(\theta)$ is equal to E of that is log likelihood function E of $\frac{\partial L}{\partial \theta}$ whole Square. This is the Fisher information statistic and we showed that this is equal to $-E$ of $\frac{\partial^2 L}{\partial \theta^2}$

square either we can take $\frac{\partial}{\partial \theta} \frac{\partial L}{\partial \theta}$ and then take the expected value or the $\frac{\partial}{\partial \theta} \frac{\partial L}{\partial \theta}$ and then take the expected value with a negative sign.

So these are the definition of Fisher information statistics and Cramer Rao theorem state that for an unbiased estimator $\hat{\theta}$ under certain regularity conditions variance of $\hat{\theta}$ is always greater than equal to $\frac{1}{I(\theta)}$ where $I(\theta)$ is the Fisher information statistic so that way it gives a bound on the variance how much smaller variance we can obtain.

(Refer Slide Time: 47:00)

MVUE through CRLB

- ❖ 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 $\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} = I(\theta)(\hat{g}(\theta) - g(\theta))$
- ❖ For an unbiased estimator

Then this here will be reached if and only if this partial derivative of log likelihood function can be written as a product of two terms $I(\theta)(\hat{\theta} - \theta)$ so this is the information statistic into $\hat{\theta} - \theta$ if the derivative partial derivative of the log likelihood function can be factorized like this then CRLB will be reached that means variance of $\hat{\theta}$, $\hat{\theta}$ will be equal to $\frac{1}{I(\theta)}$ that is the under this condition.

And also we can define the CRLB in the case of a function so variance of $\hat{g}(\theta)$ where X $\hat{g}(\theta)$ is an estimator for $g(\theta)$ is always bound by the $\frac{(g'(\theta))^2}{I(\theta)}$ this is the case when instead of a parameter we have a function. And in this case also CRLB is achieved if $\frac{\partial L}{\partial \theta}$ is equal to $I(\theta)(\hat{g}(\theta) - g(\theta))$ so this is the factorization relation and if it is satisfied then this CRLB is achieved and in that case the corresponding estimator will be MVUE because its variance is minimum.

(Refer Slide Time: 48:49)

CRLB for estimators of vector parameters

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

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

- ❖ Then the log-likelihood function is given by

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

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

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

- ❖ The Fisher Information matrix is given by

$$\mathbf{I}(\boldsymbol{\theta}) = E \left[\frac{\partial}{\partial \boldsymbol{\theta}} L(\mathbf{x}; \boldsymbol{\theta}) \left(\frac{\partial}{\partial \boldsymbol{\theta}} L(\mathbf{x}; \boldsymbol{\theta}) \right)^T \right]$$

Then we discussed this CRLB for multiple parameter or vector parameter case. So, suppose $\theta_1, \theta_2, \dots, \theta_k$ are k parameters which are represented as this vector and which characterizes the likelihood function. So, likelihood function is characterized by this set of parameters. Then log likelihood function will be so this is the likelihood function they this is the joint PDF, we are considering PDF only.

This is the joint PDF as a function of $\theta_1, \theta_2, \dots, \theta_k$ then if we take the log likelihood function then $L(\mathbf{x}; \boldsymbol{\theta})$ is equal to log of f that is the log likelihood function. And now we can take the first order partial derivative which respect to all parameters $\theta_1, \theta_2, \dots, \theta_k$ and we present it as a column vector. So, that is the partial derivative of the log likelihood function with respect to $\boldsymbol{\theta}$ that is a column vector.

And the Fisher information matrix is now given by $\mathbf{I}(\boldsymbol{\theta})$ this is a matrix now that is equal to $E \left[\frac{\partial L(\mathbf{x}; \boldsymbol{\theta})}{\partial \boldsymbol{\theta}} \left(\frac{\partial L(\mathbf{x}; \boldsymbol{\theta})}{\partial \boldsymbol{\theta}} \right)^T \right]$ that is a column vector multiplied by $\left(\frac{\partial L(\mathbf{x}; \boldsymbol{\theta})}{\partial \boldsymbol{\theta}} \right)^T$ that is a row vector. So, if we multiply this column vector by this row vector we get a matrix and then we take the expected value of individual element to get this in Fisher information matrix.

(Refer Slide Time: 50:40)

CR theorem for vector case

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

❖

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

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

$$C_{\hat{\theta}} - \mathbf{I}^{-1}(\theta) \geq 0$$

Now CR theorem for vector cases $\mathbf{I}(\theta)$ we will find out the $\mathbf{I}(\theta)$ metrics that is information matrix Fisher information matrix by this relationship this is the $\mathbf{I}(\theta)$ matrix and this can be also shown to be equal to $-E$ of $\frac{\partial}{\partial \theta} \frac{\partial}{\partial \theta} L$ this can be also shown to be equal to $\mathbf{I}(\theta)$ is equal to $\mathbf{I}(\theta)$ matrix is equal to $-E$ of $\frac{\partial}{\partial \theta} \frac{\partial}{\partial \theta} L$ vector of $\frac{\partial}{\partial \theta} L$ del theta vector. So, that way this will get as a matrix like this its component is this second order partial derivative.

So this is the second order partial derivative matrix and then we take the expectation of the individual elements with a negative sign that is the Fisher information matrix. Now assume that the likelihood function that is $f(\mathbf{x}; \theta)$ satisfy the regularity condition certain regularity conditions that we have stated in the lecture. Then the covariance matrix $E(\hat{\theta})$ or when the unbiased estimator $\hat{\theta}$ vector satisfy this relationship that is $C_{\hat{\theta}} - \mathbf{I}^{-1}(\theta) \geq 0$ $C_{\hat{\theta}}$ matrix – inverse of matrix is always greater than equal to 0 $C_{\hat{\theta}}$ matrix – inverse of θ matrices always greater than equal to 0 matrix.

What does it mean? Is that did $C_{\hat{\theta}} - \mathbf{I}^{-1}(\theta)$ that must be a positive semi definite matrix this is a positive semi definite matrix. This is the semi definite matrix. So, this way we defined the CRLB for a vector parameter case.

(Refer Slide Time: 53:07)

MVUE through sufficient statistic

❖ A statistic $T(X_1, X_2, \dots, X_N)$ of θ is called sufficient if the conditional PDF $f(x_1, x_2, \dots, x_N; \theta | T=t)$ or the conditional PMF $p(x_1, x_2, \dots, x_N; \theta | T=t)$ does not involve θ .

❖ Factorization theorem-

For continuous RVS $X_1, X_2, \dots, X_N$, the statistic $T(X_1, X_2, \dots, X_N)$ is a sufficient statistic for θ if and only if

$$f(x_1, x_2, \dots, x_N; \theta) = g(\theta, T(\mathbf{x}))h(\mathbf{x})$$

For the discrete case,

$T(\mathbf{x})$ is sufficient if and only if

$$p(\mathbf{x}; \theta) = g(\theta, T(\mathbf{x}))h(\mathbf{x})$$

Then we derived MVUE through sufficient statistic so if statistics $T(X_1, X_2, \dots, X_N)$ of θ is called a sufficient statistic if the conditional PDF that is the PDF of $x_1, x_2, \dots, x_N$ as a function of θ given T is equal to T is equal to small t this is the value of this statistic or the conditional PMF p of $x_1, x_2, \dots, x_N$ as a function of θ given T is equal to small t does not depend on θ does not involve any θ term then we say that this statistic is sufficient statistic that with information contain indeed statistic is sufficient to estimate the unknown parameter θ .

Then we established one important theorem what is known as a factorization theorem for a continuous random variable case for continuous random variables $X_1, X_2, \dots, X_N$ this statistic is a sufficient this statistics $T(X_1, X_2, \dots, X_N)$ is a sufficient statistic for θ if and only if this likelihood function is a product of two terms one is $g(\theta, T(\mathbf{x}))$ into $h(\mathbf{x})$ product of two terms one is $g(\theta, T(\mathbf{x}))$ which is a function of θ and $T(\mathbf{x})$ other term is $h(\mathbf{x})$ which is a function of $\mathbf{x}$ only it does not involve θ or $T(\mathbf{x})$.

So for the discrete case similarly we have $T(\mathbf{x})$ is sufficient if $P(\mathbf{x}; \theta)$ that is the probability mass function is a function of this quantity $g(\theta, T(\mathbf{x}))$ into $h(\mathbf{x})$ where $h(\mathbf{x})$ is independent of any θ term or $T(\mathbf{x})$ term. So, that way we defined the we derived factorization theorem to test the sufficient statistic.

(Refer Slide Time: 55:21)

MVUE through sufficient statistic

❖ Rao Blackwell theorem- Given an unbiased estimator $\hat{\theta}$, the sufficient statistic $T(\mathbf{X})$ helps us to find a better estimator $\hat{\theta} = E(\hat{\theta} / T(\mathbf{X}))$ with $\text{var}(\hat{\theta}) \leq \text{var}(\hat{\theta})$.

❖ A statistic $T(\mathbf{X})$ is said to be complete if for any bounded function $g(T(\mathbf{X}))$, the condition

$$E(g(T(\mathbf{X}))) = 0 \text{ for } \forall \theta$$

implies that

$$P(g(T(\mathbf{X})) = 0) = 1 \text{ for } \forall \theta$$

❖ If $T(\mathbf{X})$ is a complete statistic then there is only one function $g(T(\mathbf{X}))$ which is unbiased.

❖ Lehmann Scheffe theorem :

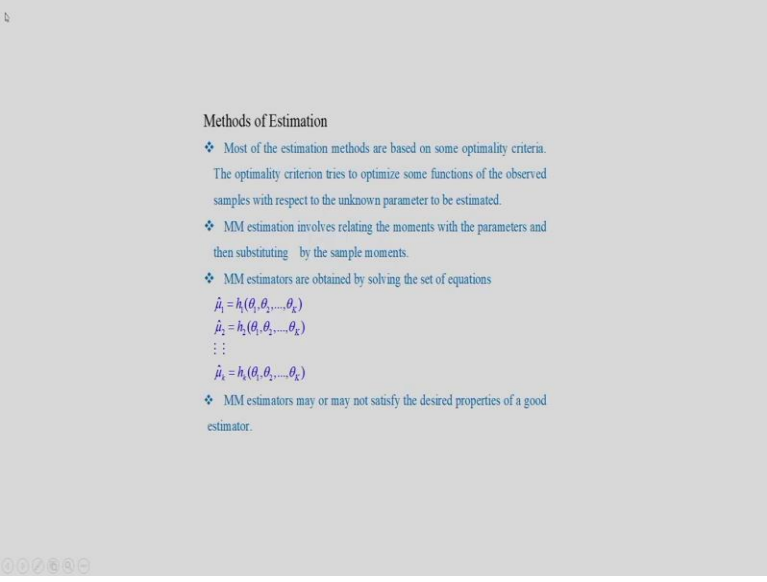
Suppose $T(\mathbf{X})$ is complete sufficient statistic for θ and $g(T(\mathbf{X}))$ is unbiased estimator based on $T(\mathbf{X})$. Then $g(T(\mathbf{X}))$ is the MVUE

Then we discussed the Rao Blackwell theorem given an unbiased estimator $\hat{\theta}$ of θ , this sufficient statistic $T(\mathbf{X})$ helps us to find a better estimator $\hat{\theta} = E(\hat{\theta} / T(\mathbf{X}))$ what is that that is equal to $E(\hat{\theta} / T(\mathbf{X}))$ its variance will be less than or equal to the variance of $\hat{\theta}$ so that is the important property we will have we can get an estimator $\hat{\theta}$ which is given by this relationship and which is unbiased this $\hat{\theta}$ will be unbiased and its variance will be lower than the variance of the other unbiased estimator $\hat{\theta}$.

So this theorem shows us how to obtain an unbiased estimator for θ given a sufficient statistic. Then we define the complete statistic a statistic $T(\mathbf{X})$ is said to be complete if for any bounded function $g(T(\mathbf{X}))$ the condition $E(g(T(\mathbf{X}))) = 0$ for all θ implies that probability that $g(T(\mathbf{X})) = 0$ is equal to 1 for all θ . So, this is the condition if this expected value is 0 then corresponding probability must be equal to 1.

So this is the complete statistic if $T(\mathbf{X})$ is a complete statistic then there is only one function of $g(T(\mathbf{X}))$ which is unbiased that is also an important result and then we got the Lehmann Scheffe theorem. Suppose $T(\mathbf{X})$ is a complete sufficient statistic for θ and $g(T(\mathbf{X}))$ is unbiased estimator based on $T(\mathbf{X})$ then $g(T(\mathbf{X}))$ is an MVUE so this is an important result and that means if it is a function of a complete sufficient statistic if the unbiased estimator is a function of a complete sufficient statistic then it will be MVUE. So that way we see how we can get MVUE through a complete sufficient statistic.

(Refer Slide Time: 57:50)



Methods of Estimation

- ❖ Most of the estimation methods are based on some optimality criteria. The optimality criterion tries to optimize some functions of the observed samples with respect to the unknown parameter to be estimated.
- ❖ MM estimation involves relating the moments with the parameters and then substituting by the sample moments.
- ❖ MM estimators are obtained by solving the set of equations
$$\hat{\mu}_1 = h_1(\theta_1, \theta_2, \dots, \theta_k)$$
$$\hat{\mu}_2 = h_2(\theta_1, \theta_2, \dots, \theta_k)$$
$$\vdots$$
$$\hat{\mu}_k = h_k(\theta_1, \theta_2, \dots, \theta_k)$$
- ❖ MM estimators may or may not satisfy the desired properties of a good estimator.

Then we discuss a few methods of estimation that means method of moments and method of maximum likelihood estimation and Bayesian estimation. So, in MM estimator method of moment estimation involves relating the moments with the parameters and then substituting by the sample moments. So, what we do we relate the moments of the distribution with the parameters and then in those relationship we substitute the moments by the corresponding sample moment and then we estimate the parameter.

So that way suppose μ_1 hat is a function of h_1 theta 1 theta 2 upto theta k, μ_2 hat is a function of h_2 theta 1 theta 2 upto theta k like that we can determine μ_k hat and they are related by this relationship therefore we can find out theta 1 hat beta 2 hat upto theta k hat by solving this set of equation. MM estimators may or may not satisfy the desired properties of good estimators. So they; but it is a very simple estimator.

(Refer Slide Time: 59:20)

Maximum likelihood estimator

The maximum likelihood estimator $\hat{\theta}_{MLE}$ is such an estimator that

$$f(x_1, x_2, \dots, x_n; \hat{\theta}_{MLE}) \geq f(x_1, x_2, \dots, x_n; \theta), \forall \theta.$$

❖ If the likelihood function is differentiable with respect to θ , then $\hat{\theta}_{MLE}$ is given

by
$$\frac{\partial f(x, \theta)}{\partial \theta} \Big|_{\hat{\theta}_{MLE}} = 0$$

or equivalently
$$\frac{\partial L(x, \theta)}{\partial \theta} \Big|_{\hat{\theta}_{MLE}} = 0$$

❖ If we have k unknown parameters $\theta = [\theta_1, \theta_2, \dots, \theta_k]^T$

then the MLEs are given by a set of equations:

$$\frac{\partial L(x, \theta)}{\partial \theta_1} \Big|_{\hat{\theta}_{MLE}} = \frac{\partial L(x, \theta)}{\partial \theta_2} \Big|_{\hat{\theta}_{MLE}} = \dots = \frac{\partial L(x, \theta)}{\partial \theta_k} \Big|_{\hat{\theta}_{MLE}} = 0$$

Then we discuss the maximum likelihood estimator, the maximum likelihood estimator $\hat{\theta}_{MLE}$ is such an estimator that the likelihood function as a function of $\hat{\theta}_{MLE}$ is greater than equal to the likelihood function as a function of θ for any θ . So, for all θ will condition consider this relationship that the likelihood function of $\hat{\theta}_{MLE}$ is always greater than this estimate is the maximum likelihood estimator.

And therefore this likelihood function is maximum for $\hat{\theta}_{MLE}$ this condition can be written in terms of this derivative condition, partial derivative of $\frac{\partial f}{\partial \theta}$ at $\hat{\theta}_{MLE}$ is equal to 0. The likelihood function if the likelihood function is differentiable with respect to θ then $\hat{\theta}_{MLE}$ is given by this relationship $\frac{\partial f}{\partial \theta}$ at $\hat{\theta}_{MLE}$ is equal to 0 or in terms of log likelihood function we can write $\frac{\partial L}{\partial \theta}$ at $\hat{\theta}_{MLE}$ is equal to 0.

And similarly if we have k unknown parameters $\theta = [\theta_1, \theta_2, \dots, \theta_k]^T$ then the MLE's are obtained by this set of equation $\frac{\partial L}{\partial \theta_1}$ at $\hat{\theta}_{MLE}$ is equal to 0 $\frac{\partial L}{\partial \theta_2}$ at $\hat{\theta}_{MLE}$ is equal to 0 $\frac{\partial L}{\partial \theta_3}$ at $\hat{\theta}_{MLE}$ is equal to 0 $\frac{\partial L}{\partial \theta_4}$ at $\hat{\theta}_{MLE}$ is equal to 0 like that where we can get a set of equations.

(Refer Slide Time: 1:01:25)

Properties of ML estimators

- ❖ MLE may be biased or unbiased.
- ❖ If a sufficient statistic $T(\mathbf{x})$ exists for θ , then $\hat{\theta}_{MLE}$ is a function of $T(\mathbf{x})$.
- ❖ If the CRLB is reached, the MLE reaches it. Thus, if

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

then, $\hat{\theta} = \hat{\theta}_{MLE}$

- ❖ The MLE is asymptotically unbiased and efficient. Thus, for large N , the MLE is approximately efficient.
- ❖ $\hat{\theta}_{MLE}$ is a consistent estimator.

And we established the properties of ML estimate or MLE estimator may be biased or unbiased if a sufficient statistic $T(\mathbf{x})$ exists for θ then $\hat{\theta}_{MLE}$ is a function of that statistic that is one important property. Again if CRLB is reached if the condition partial shall be equality CRLB equality is satisfied then MLE will reach that condition. So, if CRLB is reached then MLE is efficient. So, that means in the case where Cramer Rao bound is satisfied with equality in the case maximum likelihood estimator will have the Cramer Rao lower bound for the variance.

So that therefore if $\frac{\partial L}{\partial \theta}$ is equal to $I(\theta)(\hat{\theta} - \theta)$ then $\hat{\theta}$ must be that MLE this is the condition. The MLE is asymptotically unbiased and efficient thus for large N MLE is asymptotically unbiased for large N it will satisfy the Cramer lower bound and it is asymptotically unbiased also. Then in all cases $\hat{\theta}_{MLE}$ is a consistent estimator.

(Refer Slide Time: 1:03:11)

Bayesian estimators

❖ In Bayesian estimators, the parameter θ is assumed to be random variable with the prior PDF $f(\theta)$ or a prior PMF $p(\theta)$.

❖ The estimator uses the posterior PDF obtained using the Bayes rule:

$$f(\theta | \mathbf{x}) = \frac{f(\theta)f(\mathbf{x} | \theta)}{f_{\mathbf{x}}(\mathbf{x})}$$

❖ We associate a cost function or a loss function $C(\theta - \hat{\theta})$ with the estimator $\hat{\theta}$.

❖ A Bayesian estimator solves the optimization problem

$$\underset{\text{over } \hat{\theta}}{\text{Minimize}} \bar{C}(\hat{\theta}) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} C(\theta - \hat{\theta}) f(\mathbf{x}, \theta) d\mathbf{x} d\theta$$

❖ The MMSE estimator minimizes the mean square error and is given by

$$\hat{\theta}_{\text{MMSE}} = \int_{-\infty}^{\infty} \theta f(\theta | \mathbf{x}) d\theta = E(\theta | \mathbf{x})$$

Then we discussed about Bayesian estimators in Bayesian estimators the parameter θ is assumed to be random variable so in MLE or method of moments we consider θ to be a unknown deterministic quantity but in Bayesian approach it is a random quantity. So, therefore prior PDF $f(\theta)$ or by prior PMF $P(\theta)$ must be satisfied. The estimator uses the posterior PDF obtained by this relation $f(\theta | \mathbf{x})$ this is the posterior PDF it is given by $f(\theta)$ into $f(\mathbf{x} | \theta)$ divided by $f_{\mathbf{x}}(\mathbf{x})$.

And Bayesian estimator we associate a cost function or a loss function $C(\theta - \hat{\theta})$ with the estimator $\hat{\theta}$. So, the average value of this cost function is known as the Bayesian risk we want to minimize. A Bayesian estimate or solve the optimization problem minimize average value of $\hat{\theta}$ that is the risk over $\hat{\theta}$ and this is given by this relationship. So, because it is a function of $\mathbf{x}$ and θ therefore it is integration with respect to $\mathbf{x}$ and θ .

So that way this function we have to minimize there are different cost function we saw that means square error cost function. So, if we consider the cost function at the squared error cost function we will get the minimum mean square error estimate or minimum mean square error estimator minimizes the mean square error and it is given by this mean square error is $E(\theta - \hat{\theta})^2$.

This is the mean square error and this is minimized by the MMSE estimator and this is given by this relationship $\hat{\theta}_{MMSE}$ is conditional expectation $\hat{\theta}_{MMSE}$ is the conditional expectation of θ given X . So, that way the MMSE estimator is a Bayesian estimator which minimizes this cost function E of $\theta - \hat{\theta}$ whole square and it is given by this result E of θ given X .

(Refer Slide Time: 1:05:40)

MAP estimator

- ❖ The MAP estimator minimizes the hit-or-miss cost function and is given by $\hat{\theta}_{MAP} = \arg \max_{\theta} f(\theta | \mathbf{x})$
- ❖ If $f(\theta | \mathbf{x})$ is differentiable, $\hat{\theta}_{MAP}$ is given by the MAP equations

$$\left. \frac{\partial f(\theta | \mathbf{x})}{\partial \theta} \right|_{\hat{\theta}_{MAP}} = 0 \text{ and equivalently}$$

$$\left. \frac{\partial L(\theta | \mathbf{x})}{\partial \theta} \right|_{\hat{\theta}_{MAP}} = 0$$

Applying Bayes rule, the MAP equation can be written as

$$\left. \frac{\partial \ln(f(\theta))}{\partial \theta} \right|_{\hat{\theta}_{MAP}} + \left. \frac{\partial \ln(f(\mathbf{x} | \theta))}{\partial \theta} \right|_{\hat{\theta}_{MAP}} = 0$$

Similarly we derived the map estimator maximum a posteriori probability estimator. This minimizes another cost function what is known as the hit or miss was discontent we considered. So, this is suppose the error, if error is within $-\Delta$ by 2 and $+\Delta$ by 2 then this cost is zero otherwise cost is one that is the hit or miss cost function and the map estimate or minimizes the hit or miss cost function and is given by $\hat{\theta}_{MAP}$ is $\arg \max_{\theta} f(\theta | \mathbf{x})$ this is the posterior PDF and for what values of θ this posterior PDF is maximized that is the $\hat{\theta}_{MAP}$. So that the decorative mode of the posterior PDF.

If f of θ given X is differentiable then $\hat{\theta}_{MAP}$ is given by the map equation what is the map equation and that is partial derivative of the posterity PDF θ at a $\hat{\theta}_{MAP}$ is equal to 0 or this we can write in terms of log likelihood function also $\frac{\partial L}{\partial \theta}$ θ given X $\frac{\partial L}{\partial \theta}$ at $\hat{\theta}_{MAP}$ is equal to 0. And we can apply the Bayes rule to find out f θ given X then we can show that this map equations will result into this condition partial derivative

of likelihood function at that map plus partial derivative of the posterior density at θ hat map must be equal to 0.

So this is the bay map equations in terms of the likelihood function and posterior density function, thank you.