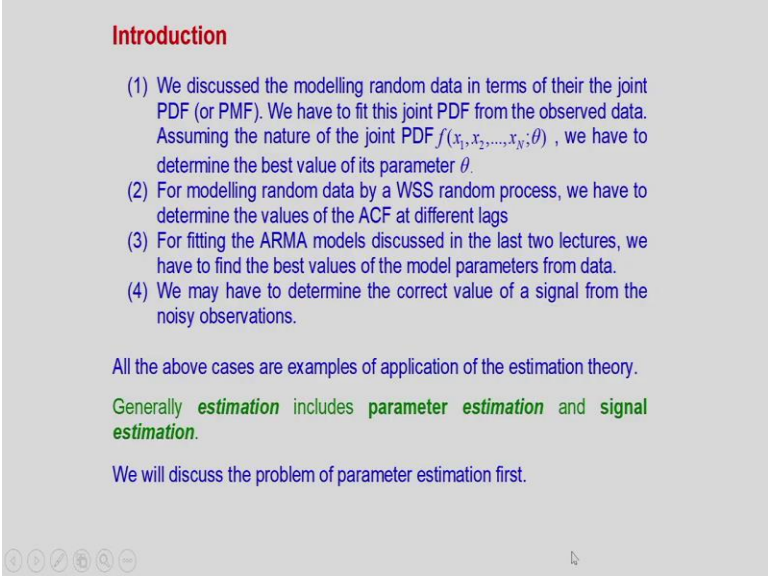


**Statistical Signal Processing**  
**Prof. Prabin Kumar Bora**  
**Department of Electronics & Electrical Engineering**  
**Indian Institute of Technology - Guwahati**

**Lecture 8**  
**Estimation Theory 1**

Hello students welcome to lecture 9 on estimation theory on 1. In this lecture, I will discuss the basics of in estimates on theory and some of the important properties of estimators. Let us start with the introduction.

**(Refer Slide Time: 00:58)**



**Introduction**

- (1) We discussed the modelling random data in terms of their the joint PDF (or PMF). We have to fit this joint PDF from the observed data. Assuming the nature of the joint PDF  $f(x_1, x_2, \dots, x_N; \theta)$ , we have to determine the best value of its parameter  $\theta$ .
- (2) For modelling random data by a WSS random process, we have to determine the values of the ACF at different lags
- (3) For fitting the ARMA models discussed in the last two lectures, we have to find the best values of the model parameters from data.
- (4) We may have to determine the correct value of a signal from the noisy observations.

All the above cases are examples of application of the estimation theory.

Generally **estimation** includes **parameter estimation** and **signal estimation**.

We will discuss the problem of parameter estimation first.

We discussed the modeling of random data in terms of their joint PDF or PMF indicates of the discrete we have to fit this joint PDF from the observed data. Assuming the nature of the joint PDF suppose we nodine f of what is the PDF and there is an unknown parameter theta, we have to determine the best value for this unknown parameter data. So, that way estimation comes for modeling random data by a WSS random process, we have to determine the value of the autocorrelation function at different lags.

Therefore, estimates are no autocorrelation function is to be done for fitting the ARMA models discussing in the last lectures, we have to find the base values of the model parameters from data. This is again estimation of the model parameters; we may have to determine the correct value of

a signal from the object noisy data. So, noisy signals from the observed noisy signals. So, that this is also a case of estimation signal estimates of signal from noisy observations.

All the above cases are examples of application of the estimation theory particularly, this is an example of signal estimation and these 1, 2 and 3 are examples of parameter estimation. And we can do also signal estimation by using parameter estimation. For example, by putting the ARMA model we can estimate the signal so that way estimation includes 2 classes parameter estimation and signal estimation. We will discuss the problem of parameter estimation.

**(Refer Slide Time: 03:03)**

**Parameter estimation**

- ❖ We have a sequence of observable random variables  $X_1, X_2, \dots, X_N$ , represented by the random vector
 
$$\mathbf{X} = \begin{bmatrix} X_1 \\ X_2 \\ \vdots \\ X_N \end{bmatrix} = [X_1 X_2 \dots X_N]^T$$
- ❖ The observed data vector  $\mathbf{x} = [x_1, x_2, \dots, x_N]^T$  is a realization or a sample of  $\mathbf{X}$
- ❖ Sometimes, we consider  $X_i$ s to be *iid-independent and identically distributed*.
- ❖  $\mathbf{X}$  is characterized by a joint PDF which depends on some unobservable parameter  $\theta$  given by
 
$$f_{\mathbf{X}}(x_1, x_2, \dots, x_N; \theta) = f_{\mathbf{X}}(\mathbf{x}; \theta) \approx f(\mathbf{x}; \theta)$$
 where  $\theta$  may be deterministic or random. Our aim is to make an inference on  $\theta$  in terms of a function  $\hat{\theta}(\mathbf{x}) = \hat{\theta}(x_1, x_2, \dots, x_N)$ .

We have a sequence of observed random variables  $X_1, X_2$  up to  $X_N$  there are  $N$  random variables represented by a random vector this is the representation it is represented the column vector and which you can write us the row vector transpose. So, that we observed data are modulus random variable. Now, particular the observed data what we observed the observed data vector small  $\mathbf{x}$  factor that is small  $x_1, x_2$  upto small  $x_N$  transpose is a realization for the samples of  $\mathbf{x}$  because random vector.

And whatever we observed at a particular time, these constitute a realization party sample up  $\mathbf{x}$ . So, small  $\mathbf{x}$  is a realization big  $\mathbf{X}$  capital  $\mathbf{X}$  is a random vector sometimes will consider exercise to be iid that is also 1 important concept independent and identically distributed all  $X_1, X_2$  and  $X_N$  are independent and each of them will have the same distribution access characterized by a

joint PDF whose depends on some unobservable parameter theta given by this we can write the joint PDF in terms of a parameter theta and this in the vector will write like this and will omit this x also. So, this will be simply f of x theta.

**(Refer Slide Time: 04:57)**

**Estimators**

- ❖ An estimator  $\hat{\theta}(X) = \hat{\theta}(X_1, X_2, \dots, X_N)$  is a rule by which we guess about the value of an unknown parameter  $\theta$ .
- ❖ It is a function of the random variables  $X_1, X_2, \dots, X_N$  and it does not involve any unknown parameters. Such a function is generally called a *statistic*.
- ❖ Being a function of random variables,  $\hat{\theta}(X)$  is random
- ❖ For a particular observation  $x_1, x_2, \dots, x_N$ , we get what is known as an estimate (not estimator)
- ❖ Thus,  $\hat{\theta}(X) = \hat{\theta}(X_1, X_2, \dots, X_N)$  is an estimator and  $\hat{\theta}(x) = \hat{\theta}(x_1, x_2, \dots, x_N)$  is an estimate.

Now, what is an estimator, an estimator  $\hat{\theta}(x) = \hat{\theta}(X_1, X_2 \text{ up to } X_N)$  and is there a rule by whose we guess about the value of an unknown parameter theta. So we will guess about the unknown parameter theta using this rule, it is a function of the random variable  $X_1, X_2 \text{ up to } X_N$ , and it does not involve any unknown parameters. This functional relationship does not involve any unknown parameter. Such a function is generally called a statistic.

So, an estimator is a statistic being a function of random variable  $\hat{\theta}(X)$  is also random. So, that way this is a random variable,  $\hat{\theta}(X)$  is a random variable for a particular observation of  $x_1, x_2 \text{ upto } x_N$ , we get what is known as an estimate, not estimator of the parameter theta thus when we consider in terms of random variable  $\hat{\theta}(x)$  that equal to  $X_1, X_2 \text{ upto } X_N$  is an estimator. And  $\hat{\theta}(x)$  of small x factor that is  $\hat{\theta}(x)$  of small  $x_1, x_2 \text{ upto } x_N$  is an estimate. So, this we have to distinguished an estimator and estimate.

**(Refer Slide Time: 06:29)**

### Example 1

Let  $X_1, X_2, \dots, X_N$  be a set of independent  $N(\mu, 1)$  random variables with the unknown mean  $\mu$ .

The joint PDF of  $X_1, X_2, \dots, X_N$  is given by

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

Then

$$\hat{\mu} = \frac{1}{N} \sum_{i=1}^N X_i \text{ is an estimator for } \mu.$$

Let us give an example 1 let us  $X_1, X_2$  upto  $X_N$  be a setup independent normal  $\mu$  random variable with unknown mean  $\mu$  parties normal distribution  $\mu$  is unknown the joint PDF of  $X_1, X_2$  upto  $X_N$  is given by this is the joint pdf and that will be productive individual marginal PDFs so, that product  $i$  going from 1 to  $N$  of  $1 / \text{root of } 2 \pi$  into  $e$  to the power  $-1/2$  of  $x_i - \mu$  whole squared.

So, using the property of exponential concern, we can write it in terms of a sum here. So, this is the joint pdf of data. And here unknown quantities  $\mu$  and if we take  $\hat{\mu} = 1 / N$  summation  $X_i$ ,  $i$  is going from 1 to  $N$  is an estimator for  $\mu$ . It does not involve any unknown parameter and it is a function of the random variables. So it is an estimator for  $\mu$ .

**(Refer Slide Time: 07:42)**

### Example 2

Suppose we have DC voltage  $\theta$  corrupted by noise  $V_i$  and the observed data  $X_i, i = 1, 2, \dots, N$  are given by

$$X_i = \theta + V_i$$

Suppose  $V_i$ s are independent and each is distributed as  $N(0, \sigma^2)$ . Thus  $X_i$ s are iid  $N(\theta, \sigma^2)$  random variables.

Then,  $\hat{\theta} = \frac{1}{N} \sum_{i=1}^N X_i$  is an estimator for  $\theta$ .

But question will arise how we can get estimator. Example 2 this is a practical example. Suppose we have a DC voltage  $\theta$  corrupted by noise  $V_i$  and the observed  $X_i, i = 1, 2$  up to  $N$  are given by  $X_i = \theta + V_i$ . Now, we assume that  $V_i$ s are independent and each is distributed as normal  $0$  sigma squared, it is  $0$  mean sigma squared, the variance thus  $X_i$ s, because  $X_i$  is the sum of  $\theta + V_i$ , this is the  $0$  mean so,  $X_i$ s are iid and it is normal.

$\theta$  sigma squared random variable so it is normal distribution with mean  $\theta$  and variance sigma squared. Then again, since it is a normal distribution, this  $\hat{\theta} = \frac{1}{N} \sum_{i=1}^N X_i$  is an estimator of  $\theta$ .

(Refer Slide Time: 08:59)

### Properties of an estimator

❖ A good estimator should satisfy some properties. These properties are described in terms of the mean and variance of the estimator.

#### Unbiased estimator

❖ An estimator  $\hat{\theta}$  of  $\theta$  is said to be *unbiased* if and only if  $E\hat{\theta} = \theta$ . Unbiasedness means that on the average the estimator gives the true value of the parameter. It is a desirable property.

❖ The quantity  $b(\hat{\theta}) = E\hat{\theta} - \theta$  is called the bias of the estimator. It is desirable that  $b(\hat{\theta})$  decreases as  $n$  increases.

❖  $\hat{\theta}$  is said to be an asymptotically unbiased if  $\lim_{n \rightarrow \infty} E\hat{\theta} = \theta$

Unbiasedness is necessary but not sufficient to make an estimator a good one.

Let us discuss the desirable properties of an estimator. A good estimator should satisfy some properties, these properties are described in terms of the mean and the variance of the estimator. A part will discuss unbiased estimator and estimator  $\hat{\theta}$  of  $\theta$  is said to be unbiased if and only  $E \text{ of } \hat{\theta} = \theta$  unbiasedness means that on the average the estimator gives the true value of the parameter.

So on the values of estimator gives the 2 value of the parameter, it is a desirable property. Now, if it is not an unbiased estimator supposed  $\hat{\theta}$  is not an unbiased estimator then  $b \text{ of } \hat{\theta} = E \text{ of } \hat{\theta} - \theta$  is called the bias of the estimator. So, this quantity difference between expected value and the true value is called a bias of the estimator, it is desirable that  $b \text{ of } \hat{\theta}$  decreases as  $n$  increases eventually go down to 0.

$\hat{\theta}$  is said to be an asymptotically unbiased estimator if  $\lim_{n \rightarrow \infty} E \hat{\theta} = \theta$ . So, as we have large number of  $\theta$ , then  $E \hat{\theta}$  will become close to  $\theta$ . unbiasedness is necessary, but not sufficient to make an estimator a good one we have to consider other properties.

**(Refer Slide Time: 10:51)**

**Example 3: Example 2 revisited**

Suppose we have DC voltage  $\theta$  corrupted by noise  $V_i$  and the observed data  $X_i, i = 1, 2, \dots, N$  are given by

$$X_i = \theta + V_i$$

If  $V_i$ s are 0-mean, then

$$EX_i = E\theta + EV_i = \theta$$

For  $\hat{\theta} = \frac{1}{N} \sum_{i=1}^N X_i$ , we get

$$E\hat{\theta} = \frac{1}{N} \sum_{i=1}^N EX_i = \frac{1}{N} \times N\theta = \theta$$

Therefore,  $\hat{\theta}$  is an unbiased estimator for  $\theta$ .

Will give another example, example 3 example 2 revisited. Suppose we have a DC voltage  $\theta$  corrupted by noise  $V_i$  and the observed data  $X_i, i = 1, 2, \dots, N$  are given by  $X_i = \theta + V_i$  if  $V_i$  is 0 mean then  $EX_i$  will be  $= E \text{ of } \theta + E \text{ of } V_i = \theta$ . For  $\hat{\theta} = \frac{1}{N} \sum_{i=1}^N X_i$

i, i is going from 1 to N that is the estimator we are considering E of theta hat will be = 1 / N into summation i going from 1 to N of E of X i. Now, E of X i = theta. Therefore, this will be = 1 / N and into N theta that will = theta, therefore this theta hat is an unbiased estimator of theta.

(Refer Slide Time: 11:56)

**Example 4**  
Suppose  $X_1, X_2, \dots, X_N$  are iid and we have two estimators for the variance  $\sigma^2$ :

$$\hat{\sigma}_1^2 = \frac{1}{N} \sum_{i=1}^N (X_i - \hat{\mu})^2 \quad \text{and} \quad \hat{\sigma}_2^2 = \frac{1}{N-1} \sum_{i=1}^N (X_i - \hat{\mu})^2$$

We can show that  $\hat{\sigma}_2^2$  is an unbiased estimator. For this

$$\begin{aligned} E \sum_{i=1}^N (X_i - \hat{\mu})^2 &= E \sum_{i=1}^N (X_i - \mu + \mu - \hat{\mu})^2 \\ &= \sum_{i=1}^N (E(X_i - \mu)^2 + E(\mu - \hat{\mu})^2 + 2E(X_i - \mu)(\mu - \hat{\mu})) \\ \therefore E \sum_{i=1}^N (X_i - \hat{\mu})^2 &= N\sigma^2 + \sum_{i=1}^N E(\mu - \hat{\mu})^2 + 2 \sum_{i=1}^N E(X_i - \mu)(\mu - \hat{\mu}) \end{aligned}$$

Now  $E(\mu - \hat{\mu})^2 = E \left( \mu - \frac{\sum X_i}{N} \right)^2$

$$\begin{aligned} &= \frac{E}{N^2} (N\mu - \sum X_i)^2 = \frac{E}{N^2} \sum (X_i - \mu)^2 \\ &= \frac{E}{N^2} \sum (X_i - \mu)^2 + \sum_{i,j=1}^N E(X_i - \mu)(X_j - \mu) \\ &= \frac{E}{N^2} \sum (X_i - \mu)^2 \quad (\text{because of independence}) \\ &= \frac{\sigma^2}{N} \end{aligned}$$

We will consider another example this is the estimator for the variance suppose  $X_1, X_2$  upto  $X_N$  are iid independent and identically distributed. And we have to estimate the part sigma squared. Sigma squared is the variance, there are 2 estimators, 1 is sigma 1 hat square that = 1 / N summation i going from 1 to N of  $X_i - \mu$  hat whole squared. So, this is 1 estimator and 2 estimator is same as the summation, but that is = 1 / N - 1 and here 1 / N.

And so that way we have 2 estimator sigma 1 hat squared and sigma 2 hat squared. We can show that this estimator sigma 2 hat squared is an unbiased estimator, but this is not an unbiased estimator for this we will consider what is your  $X_i - \mu$  hat whole squared part? Will determine that  $X_i - \mu$  hat whole squared i going from 1 to N that will write in terms of mu. So, that way  $x_i - \mu + \mu - \mu$  hat whole squared.

So, this quantity now we can expand first term we have  $X_i - \mu$  hat whole squared, the second term will be your  $\mu - \mu$  hat whole square and then crossed terms will be there twice you have  $X_i - \mu$  into  $\mu - \mu$  hat. So, I know that this is iid, so, that way they are N such terms because of the summation. So, we can write this expectation the E = first term will be N sigma squared and

second term will be summation up this quantity here  $\mu - \hat{\mu}$  whole square and third term will be summation of this quantity of  $x_i - \mu$  into  $\mu - \hat{\mu}$ .

So, let us see what this term will equal to so, this is your  $\mu - \hat{\mu}$  whole square, this expression is important for us. So, your  $\mu - \hat{\mu}$  whole squared that is billion of  $\hat{\mu} =$  your  $\mu - \sum x_i / N$  whole square. So, if I take  $N$ , then because of this squared  $E$  of, and  $\mu - \sum x_i$  whole square /  $N$  square and this we can write as your summation  $x_i - \mu$  whole square /  $N$  square.

Now, this is again a summation that summation we can write in terms of the individual squadrons  $E$  of summation of  $x_i - \mu$  whole square and then all cross terms we have to consider, because these are cross terms  $x_i - \mu$  into  $x_j - \mu$  in this form. Therefore, joint expectation of this quantities will be 0, because of independent exercise and  $x_j$  are independent therefore, this expression will become 0. So, therefore, this will be simply  $E$  of summation of  $x_i - \mu$  whole squared divided by  $N$  squared and this is equal to sigma squared by  $N$ .

And because it is identical with variance will be same for all  $x_i$  and this time will be  $= E x_i - \mu$  whole squared will be sigma square so that way we will get sigma squared here divided by  $N$  squared it will be sigma square by  $N$ . So, this is an important observation therefore, this quantity  $E$  of  $\mu - \hat{\mu}$  whole square = sigma square by  $N$  this result will be used later on also we have to remember this result.

**(Refer Slide Time: 16:12)**

Contd...

$$\text{Also } E(X_i - \mu)(\mu - \hat{\mu}) = -\frac{\sigma^2}{N}$$

$$\begin{aligned} \therefore E \sum_{i=1}^N (X_i - \hat{\mu})^2 &= N\sigma^2 + \sigma^2 - 2\sigma^2 \\ &= (N-1)\sigma^2 \\ \Rightarrow E\hat{\sigma}_1^2 &= \frac{1}{N} E \sum (X_i - \hat{\mu})^2 \\ &= \frac{N-1}{N} \sigma^2 \neq \sigma^2 \end{aligned}$$

and

$$E\hat{\sigma}_2^2 = \frac{1}{N-1} E \sum (X_i - \hat{\mu})^2 = \sigma^2$$

$\hat{\sigma}_2^2$  is an unbiased estimator of  $\sigma^2$ , while  $\hat{\sigma}_1^2$  is not.

$\hat{\sigma}_1^2$  is asymptotically unbiased.

And similarly, we can do  $X_i - \mu$  into  $\mu - \hat{\mu}$  now,  $\hat{\mu}$  to involve all the random variables, but only in the case where it is  $X_i$ , then only this expectation will become non-zero otherwise this expectation will become 0. So, that way this expression will become  $-\sigma^2$  by  $N$  and therefore, we have to determine this expression, your  $E$  of summation  $X_i - \hat{\mu}$  whole square,  $i$  going from 1 to  $N$ , this =  $N\sigma^2$  then this term plus this term.

So, we can write this is  $E$  of summation  $i$  going from 1 to  $N$  of  $X_i - \hat{\mu}$  whole squared =  $N\sigma^2 + \sigma^2 - 2\sigma^2$ . So, that way it will become because this is  $\sigma^2$  squared, this is  $-2\sigma^2$  squared. So, that way it will become  $N-1\sigma^2$  squared. So, this will imply that now, if we have to find out that  $E$  of  $\hat{\sigma}_1^2$  square that is scaling factor is  $1/N$  into this quantity. So, that way it will become  $(N-1)/N$  into  $\sigma^2$  squared.

So, this is not equal to  $\sigma^2$  so, if I consider this  $E$  of  $\hat{\sigma}_1^2$  square that is not equal to  $\sigma^2$  squared therefore, it is not an unbiased estimator. But if we consider  $\hat{\sigma}_2^2$  squared =  $1/(N-1)$  into  $E$  of summation  $X_i - \hat{\mu}$  whole square =  $\sigma^2$  squared therefore  $\hat{\sigma}_2^2$  whole square that will be unbiased estimator. So  $\hat{\sigma}_2^2$  squared is an unbiased estimator of  $\sigma^2$  while  $\hat{\sigma}_1^2$  squared is not.

That is why in determining the sample variance we always divide by  $N - 1$  to make it an unbiased estimator. But as  $N$  tends to infinity this quantity will become 1, therefore  $\sigma^2$  is asymptotically unbiased.

(Refer Slide Time: 18:48)

**Variance of the estimator**

- ❖ The variance of the estimator  $\hat{\theta}$  is given by
 
$$\text{var}(\hat{\theta}) = E(\hat{\theta} - E(\hat{\theta}))^2$$
- ❖ For the unbiased case
 
$$\text{var}(\hat{\theta}) = E(\hat{\theta} - \theta)^2$$

The variance of the estimator should be as low as possible.
- ❖ An unbiased estimator  $\hat{\theta}$  is called a *minimum variance unbiased estimator (MVUE)* if
 
$$E(\hat{\theta} - \theta)^2 \leq E(\hat{\theta}' - \theta)^2$$

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

So we consider the main of the estimator will consider the variance of the estimator. The variance of the estimator  $\hat{\theta}$  is given by variance of  $\hat{\theta} = E(\hat{\theta} - E(\hat{\theta}))^2$ . That is the variance this is the random variable and it is mean. So, that way this is the variance for the unbiased case this variance of the  $\hat{\theta}$  is will be simply of  $\hat{\theta} - \theta$  whole square because  $E(\hat{\theta}) = \theta$ .

The variance of an estimator should be as low as possible and unbiased estimator  $\hat{\theta}$  is called a minimum variance unbiased estimator MVUE if  $E(\hat{\theta} - \theta)^2 \leq E(\hat{\theta}' - \theta)^2$  that is variance of  $\hat{\theta}$  is less than equal to variance of  $\hat{\theta}'$  that is variance of  $\hat{\theta}$  thus so variance of  $\hat{\theta}$  is less than equal to variance of  $\hat{\theta}'$  thus were  $\hat{\theta}$  thus is any other unbiased estimator.

(Refer Slide Time: 20:00)

### Mean square error of the estimator (MSE)

- ❖  $MSE = E(\hat{\theta} - \theta)^2$  is an important estimation criterion.
- ❖ MSE should be as small as possible. Out of all unbiased estimator, the MVUE has the minimum mean square error.
- ❖ MSE is related to the bias and variance as shown below.

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

$$\begin{aligned} MSE &= E(\hat{\theta} - \theta)^2 = E(\hat{\theta} - E\hat{\theta} + E\hat{\theta} - \theta)^2 \\ &= E(\hat{\theta} - E\hat{\theta})^2 + E(E\hat{\theta} - \theta)^2 + 2E(\hat{\theta} - E\hat{\theta})(E\hat{\theta} - \theta) \\ &= E(\hat{\theta} - E\hat{\theta})^2 + E(E\hat{\theta} - \theta)^2 + 2(E\hat{\theta} - E\hat{\theta})(E\hat{\theta} - \theta) \\ &= \text{var}(\hat{\theta}) + b^2(\hat{\theta}) + 0 \end{aligned}$$

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

So we discussed mean variance another term mean square error MSE.  $MSE = \text{theta} - \text{theta hat}$  whole square, minimizing the MSE is an important estimation criterion. MSE should be as small as possible out of all unbiased estimators the MVUE has the minimum mean square error so in the case of unbiased estimator. MVUE will have the list MSE now, MSE is related to bias and variance as shown below  $MSE = \text{variance of theta hat} + \text{bias of theta hat squared}$ .

This you can probably  $MSE = E \text{ of theta hat} - \text{theta}$  whole square by definition that is equal to now we can write  $E \text{ of theta hat} - E \text{ of theta hat} + E \text{ of theta hat} - \text{theta}$  square we have subtracting  $E \text{ of theta hat}$  and adding  $E \text{ of theta hat}$ . Now we expand it  $E \text{ of theta hat} - E \text{ theta hat}$  whole square  $E \text{ of } E \text{ theta hat} - \text{theta}$  whole square + twice a term  $2 \text{ of theta hat} - E \text{ theta hat}$  into  $E \text{ of theta hat} - \text{theta}$ . Now, this term, first term is the variance of theta hat.

And second time, because it is a constant quantity it will remain as  $E \text{ of theta hat} - \text{theta}$  only. And therefore, this quantity will be biased of theta hat squared and the third term because  $E \text{ of theta hat} - E \text{ of } E \text{ of theta hat}$  that will be same as the  $E \text{ of theta hat}$  because this is a constant quantities or this term will become 0 therefore, this will become 0. So therefore what we get  $MSE = \text{variance of theta hat} + \text{biased of theta hat}$ . So this is 1 important relationship and when this quantity will become 0, bias 0, and in this case  $MSE = \text{variance of theta hat}$ .

**(Refer Slide Time: 22:19)**

### Consistent Estimators

- ❖ As we have more data, the quality of estimation should be better. This idea is used in defining the consistent estimator.
- ❖ An estimator  $\hat{\theta}$  is called a consistent estimator of  $\theta$  if  $\hat{\theta}$  converges in probability to  $\theta$ .  

$$\lim_{N \rightarrow \infty} P(|\hat{\theta} - \theta| \geq \varepsilon) = 0 \text{ for any } \varepsilon > 0$$
- ❖ Thus a consistent estimator converges to the true value in probability.
- ❖ Less rigorous test is obtained by applying the Chebyshev Inequality

$$P(|\hat{\theta} - \theta| \geq \varepsilon) \leq \frac{E(\hat{\theta} - \theta)^2}{\varepsilon^2} \quad \text{MSE}$$

Therefore, if

$$\lim_{N \rightarrow \infty} E(\hat{\theta} - \theta)^2 = 0 \text{ then } \hat{\theta} \text{ will be a consistent estimator.}$$

Now we will discuss consistent estimators as we have more data the quality of estimations will be better. This idea is used in defining the consistent estimator. So this estimator is a good estimator, if we have large amount of data and estimator,  $\theta$  hat is called a consistent estimator if  $\theta$  hat converges in probability to  $\theta$ . So converges in probability that is defined in this way. So this is the limit as  $N$  tends to infinity of the probability of mod of  $\theta$  minus  $\theta$  greater than or equal to  $\varepsilon$ .

So, that means probability of the deviation as  $N$  tends to infinity that would go down to 0 for any  $\varepsilon$  greater than 0. So that we have defined a consistent estimator thus, a consistent estimator converges to the true value of data in probability less rigorous this because here we have to determine the probability. So they are less rigorous test is obtained by applying this Chebyshev inequality.

Now probability of this derivation is less than equal to  $E$  of  $\theta$  hat -  $\theta$  hat whole squared, that is nothing but the MSE divided by  $\varepsilon$  square. So therefore, if this quantity will be 0, then this will also become 0. Therefore if limit of  $E$  of  $\theta$  hat -  $\theta$  whole square as  $N$  tends to infinity = 0, then  $\theta$  hat will be a consistent estimator. So this is the consistent estimator, this is the test for consistent estimator will be using.

**(Refer Slide Time: 24:15)**

### Consistent estimator...

- ❖ If  $\hat{\theta}$  is an unbiased estimator of  $\theta$ , then

$$MSE = \text{var}(\hat{\theta})$$

Therefore, if the estimator  $\hat{\theta}$  is unbiased and  $\text{var}(\hat{\theta}) \rightarrow 0$ ,  
as  $N \rightarrow \infty$  then  $\hat{\theta}$  will be a consistent estimator.

- ❖ Note that consistency is an asymptotic property. As we get more and more data, the estimates become better and better.

Now if  $\theta$  is unbiased in the case MSE is same as variance of  $\theta$  does if the estimator  $\theta$  is unbiased and variance of  $\theta$  goes down to 0 as  $N$  tends to infinity, then  $\theta$  will be a consistent estimator. Note that consistency is an asymptotic property. As we get more and more data the estimates become better and better that is the concept behind consistent estimator.

(Refer Slide Time: 24:48)

### Example 5

Suppose  $X_1, X_2, \dots, X_N$  are iid with unknown  $\mu$  and known variance

$$\sigma^2. \text{ Let } \hat{\mu} = \frac{1}{N} \sum_{i=1}^N X_i$$

$$\text{Then, } E\hat{\mu} = \frac{1}{N} \sum_{i=1}^N EX_i = \mu \text{ so that } \hat{\mu} \text{ is unbiased.}$$

Also

$$\begin{aligned} \text{var}(\hat{\mu}) &= E(\hat{\mu} - \mu)^2 \\ &= E\left(\frac{\sum X_i}{N} - \mu\right)^2 \\ &= \frac{\sigma^2}{N} \end{aligned}$$

$$\text{Clearly } \lim_{N \rightarrow \infty} \text{var}(\hat{\mu}) = \lim_{N \rightarrow \infty} \frac{\sigma^2}{N} = 0$$

Therefore  $\hat{\mu}$  is a consistent estimator of  $\mu$ .

Let us consider 1 example  $X_1$  upto  $X_N$  are iid with unknown  $\mu$  known variance  $\sigma^2$  and suppose  $\hat{\mu} = 1/N \sum_{i=1}^N X_i$ . So,  $\hat{\mu}$  will be  $= \sum_{i=1}^N X_i / N$ . So, that way it will be  $\mu$  so that  $\hat{\mu}$  is unbiased and variance of  $\hat{\mu}$  is unbiased. Variance of  $\hat{\mu}$  now, this variance we have that equal to variance of  $\hat{\mu}$  =

E of  $\mu$  hat -  $\mu$  whole squared that we have already determined in our example 4. So, that way, this will be equal to a simply  $\sigma^2 / N$  so, variance of  $\mu$  hat =  $\sigma^2 / N$  because  $X_i$  that iid.

So, as limit  $N$  tends to infinity, this variance will become 0. So, limit of variance of  $\mu$  hat as  $N$  tends to infinity that = limit of  $\sigma^2 / N$ ,  $N$  tends to infinity that is 0 therefore,  $\mu$  hat is a consistent estimator  $\mu$ . So, that is sample mean is not only unbiased, but it is also consistent. That means, if we have more and more a number of data, then variants will go down to the 0 and we will get the true value of  $\mu$ . So, that is the idea behind consistent estimators.

(Refer Slide Time: 26:32)

**Efficient Estimator**

❖ Suppose  $\hat{\theta}_1$  and  $\hat{\theta}_2$  be two unbiased estimator of  $\theta$  with  $\text{var}(\hat{\theta}_1) < \text{var}(\hat{\theta}_2)$ . The relative efficiency of the estimator  $\hat{\theta}_2$  with respect to the estimator  $\hat{\theta}_1$  is defined by

$$\text{Relative Efficiency} = \frac{\text{var}(\hat{\theta}_1)}{\text{var}(\hat{\theta}_2)}$$

Particularly, if  $\hat{\theta}_1$  is an MVUE, it is called an *efficient estimator* and the *absolute efficiency* of an unbiased estimator with respect to this estimator.

I will introduce another term efficient estimator supposed  $\theta_1$  hat and  $\theta_2$  hat are 2 unbiased estimator of  $\theta$ , with variance of  $\theta_1$  hat < variance of  $\theta_2$  hat. This is unbiased estimator. The relative efficiency of estimator of  $\theta_2$  hat with respect to estimator  $\theta_1$  hat is defined by this ratio. Variance of  $\theta_1$  hat divided by variance of  $\theta_2$  hat this variance is less therefore, this number will be less than 1 particularly if  $\theta_1$  hat is an MVUE.

Then  $\theta_1$  hat will be called an efficient estimator and the absolute efficiency of an unbiased estimator which respect to this estimator. So, this is a relative efficiency, but absolute efficiency of an unbiased estimator will be determined which respect to the MVUE. So, here and this will

be the estimator with minimum variance. Then we will call this efficiency at the absolute efficiency.

**(Refer Slide Time: 27:46)**

**Example**

Suppose  $X_1, X_2, \dots, X_N$  are i.i.d. normal random variables with unknown mean  $\mu$  and  $\hat{\mu}$  and  $\hat{\mu}_1$  are two estimators of  $\mu$  given by

$$\hat{\mu} = \frac{1}{N} \sum_{i=1}^N X_i \text{ and } \hat{\mu}_1 = \frac{1}{2}(X_1 + X_N)$$

We have shown that  $\text{var}(\hat{\mu}) = \frac{\sigma^2}{N}$ .

It can be shown that  $\text{var}(\hat{\mu}_1) = \frac{\sigma^2}{2}$ .

Therefore,  $\text{Efficiency of } (\hat{\mu}_1) = \frac{2}{N}$

Now will give another example, suppose  $X_1, X_2$  upto  $X_N$  are iid random variables with unknown  $\mu$  and  $\hat{\mu}$  and  $\hat{\mu}_1$  are 2 estimators of  $\mu$  is given by  $\hat{\mu}$  this is the sample mean as it was well, and  $\hat{\mu}_1$  suppose we defined that equal to half of  $X_1 + X_N$ . And now both are unbiased estimator and it see that variance up  $\hat{\mu}$ , this quantity sample means is equal to  $\sigma^2 / N$ . And we can do that by using the same formula  $N = 2$  variance of  $\hat{\mu}_1$  will be  $= \sigma^2 / 2$ .

Therefore, efficiency relative efficiency will  $\hat{\mu}_1$  will be  $2 / N$  if I consider of this quantity divided by this quantity, then we will get  $2 / N$  so, this is the relative efficiency of estimator  $\hat{\mu}_1$  which respects to  $\hat{\mu}$ .

**(Refer Slide Time: 28:54)**

### Summary

- ❖ The observed random data  $X_1, X_2, \dots, X_N$  are characterized by a joint PDF which depends on some unobservable parameter  $\theta$ .
- ❖ An estimator  $\hat{\theta}(X) = \hat{\theta}(X_1, X_2, \dots, X_N)$  is a rule by which we guess about the value of the unknown parameter  $\theta$ .
- ❖ An estimator  $\hat{\theta}$  of  $\theta$  is said to be *unbiased* if and only if  $E\hat{\theta} = \theta$ . Unbiasedness means that on the average the estimator gives the true value of the parameter. It is a desirable property.
- ❖  $\hat{\theta}$  is said to be an asymptotically unbiased if  $\lim_{N \rightarrow \infty} E\hat{\theta} = \theta$ .
- ❖ An unbiased estimator  $\hat{\theta}$  is called a *minimum variance unbiased estimator (MVUE)* if
$$\text{var}(\hat{\theta}) \leq \text{var}(\hat{\theta}')$$
where  $\hat{\theta}'$  is any other unbiased estimator.

Let us summarize the observed random data  $X_1, X_2$  upto  $X_N$  these are the random variables are modeled as random variables and are characterized by joint PDF which depends on some unobservable parameter  $\theta$ . An estimator  $\hat{\theta}(X) = \hat{\theta}(X_1, X_2 \text{ upto } X_N)$  is a rule by which we guess about the value of the unknown parameter  $\theta$  and estimator  $\hat{\theta}$  of  $\theta$  is said to be unbiased, we introduce what is an unbiased estimator.

So unbiased if and only if  $E\hat{\theta} = \theta$ . Unbiasedness means that on the average the estimator gives the true value of the parameter, it is a desirable property.  $\hat{\theta}$  is said to be asymptotically unbiased, if limit of  $E\hat{\theta}$  as  $N$  tends to infinity  $= \theta$ . Then we discussed about MVUE an unbiased estimator  $\hat{\theta}$  is called a minimum variance unbiased estimator (MVUE) if variance of  $\hat{\theta}$  is less than or equal to variance of  $\hat{\theta}'$  thus so, where  $\hat{\theta}'$  is any other unbiased estimator. So, in the case of MVUE variance is least minimum variance unbiased estimator.

**(Refer Slide Time: 30:33)**

### Summary ...

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

$$MSE = E(\theta - \hat{\theta})^2$$

- ❖ MSE is related to the variance and bias

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

- ❖ An estimator  $\hat{\theta}$  is called a consistent estimator of  $\theta$  if  $\hat{\theta}$  converges in probability to  $\theta$ . Thus,

$$\lim_{N \rightarrow \infty} P(|\hat{\theta} - \theta| \geq \varepsilon) = 0 \text{ for any } \varepsilon > 0$$

- ❖ If the estimator  $\hat{\theta}$  is unbiased and  $\text{var}(\hat{\theta}) \rightarrow 0$ , as  $N \rightarrow \infty$  then

$\hat{\theta}$  will be a consistent estimator.

- ❖ The relative efficiency of an estimator  $\hat{\theta}_2$  with respect to  $\hat{\theta}_1$  is given by

$$\text{Relative Efficiency} = \frac{\text{var}(\hat{\theta}_1)}{\text{var}(\hat{\theta}_2)}$$

We also discuss about the mean square error is MSE of an estimator given by  $MSE = E(\theta - \hat{\theta})^2$ . MSE is related to the variance and bias by this relationship  $MSE = \text{var}(\hat{\theta}) + b^2(\hat{\theta})$ . Then we define a consistent estimator and estimator  $\hat{\theta}$  is called a consistent estimator of  $\theta$  if  $\hat{\theta}$  converges in probability to  $\theta$ , thus, what we have limit and  $N$  turns to infinity of probability of derivation  $|\hat{\theta} - \theta| \geq \varepsilon = 0$  for any  $\varepsilon > 0$ .

So for any  $\varepsilon > 0$  if a probability of this derivation goes down to 0 as  $N$  tends to infinity, then the estimator will be consistent. If the estimator  $\hat{\theta}$  is unbiased, and variance of  $\hat{\theta}$  goes down to 0 as  $N$  tends to infinity, then  $\hat{\theta}$  will be a consistent estimator. This is a test for the consistencies of an unbiased estimator variance of  $\hat{\theta}$  go down to 0 as  $N$  tends to infinity. The relative efficiency of an estimator  $\hat{\theta}_2$  with respect to  $\hat{\theta}_1$ , which has lower variance is given by relative efficiency is equal the variance of  $\hat{\theta}_1$  hat divided by variance of  $\hat{\theta}_2$  hat. Thank You.