

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

Lecture 13
Method of Moments and Maximum Likelihood Estimators

Hello students welcome to lecture 13 on methods of moments and maximum likelihood estimators.

(Refer Slide Time: 00:55)

General Estimation Techniques ...

- ❖ Most of these methods are based on some optimality criteria. An optimality criterion tries to optimize some functions of the random samples with respect to the unknown parameter to be estimated.
- ❖ Some of the most popular estimation techniques are:
 - Method of moments
 - Maximum likelihood method
 - Bayesian methods.
 - Least squares methods

We will discuss the first three techniques in this module. The least squares principle will be discussed in a later module.

In this lecture we will introduce some general methods for parameter estimation. We saw that MVUE minimum variance unbiased estimator is the most desirable estimator. We discussed two approaches to find the MVUE through CRLB and MVUE through a complete sufficient statistic. The above approaches may not be feasible for practical models of random data. In a practical situation, we have to apply some general techniques to construct a good estimator.

The same technique applied on different probability models may produce different estimation rules. The goodness of an estimator is measured in terms of the desired properties of an estimator like unbiasedness, consistence and efficiency. Most of these methods are based on some optimality criteria. An optimality criterion tries to optimize some functions of the random sample with respect to the unknown parameter to be estimated.

Some of the popular estimation techniques are method of moments, maximum likelihood method, Bayesian methods, least squares methods. We will discuss the first three techniques in this module. We will discuss method of moments, maximum likelihood method and

Bayesian methods. The least squares principle will be discussed in a later module. When we discuss about adaptive filters the time we will introduce least squares estimation method.

(Refer Slide Time: 02:52)

Method of Moments

- ❖ The method of moments (MM) is a simple criterion for parameter estimation. When other methods are mathematically intractable, an MM estimator is a simple alternative.
- ❖ Suppose $X_1, X_2, \dots, X_N$ are iid random samples with the joint probability density function $f(x_1, x_2, \dots, x_N; \theta_1, \theta_2, \dots, \theta_K)$ which depends on unknown parameters $\theta_1, \theta_2, \dots, \theta_K$.
- ❖ The r -th moment of each X_i is given by EX_i^r $r = 1, 2, \dots$ Thus,

For $r = 1$ $EX_1 = \text{Mean of } X_1$
 For $r = 2$ $EX_1^2 = \text{Mean-square value of } X_1$
 and so on

We will start with method of moments; the method of moments abbreviated as MM is a simple criterion for parameter estimation, it is the simplest criterion. When other methods are mathematically intractable, an MM estimator is a simple alternative. First we will discuss what are moments, suppose X_1, X_2 up to X_N are iid random samples with the joint probability density function $f(x_1, x_2$ up to x_N as a function of θ_1, θ_2 up to θ_K .

This density function depends on unknown parameters θ_1, θ_2 up to θ_K , now the r -th moment of X_i is given by E of X_1 to the power r , because they are iid any random variable we can consider E of X_1 to the power r that is the r -th moment and we can define for r is equal to 1, 2 etcetera any positive integer the loop. So for r is equal to 1, we have E of X_1 that is the mean of X_1 .

Similarly for r is equal to 2 we have E of X_1 square that is the mean square value of X_1 and so on. So that way we can define the moment of a random variable X_i

(Refer Slide Time: 04:21)

Sample Moments

❖ The sample moments are given by

$$\hat{\mu}_r = \frac{1}{N} \sum_{i=1}^N x_i^r, \quad r=1,2,\dots$$

For example,

$$\hat{\mu}_1 = \hat{\mu} = \frac{1}{N} \sum_{i=1}^N x_i \quad \text{and} \quad \hat{\mu}_2 = \frac{1}{N} \sum_{i=1}^N x_i^2$$

❖ Note that

$$E \hat{\mu}_r = \frac{1}{N} \sum_{i=1}^N E x_i^r = E X_1^r \quad \text{and}$$

$$\text{var}(\hat{\mu}_r) = \frac{1}{N^2} \sum_{i=1}^N \text{var} X_i^r = \frac{\text{var} X_1^r}{N}$$

$$\text{So that } \lim_{N \rightarrow \infty} \text{var}(\hat{\mu}_r) = \lim_{N \rightarrow \infty} \frac{\text{var} X_1^r}{N} = 0$$

$\therefore \hat{\mu}_r$ is an unbiased and a consistent estimator.

Now sample moments are the moments estimated from data. The sample moments are given by suppose r-th moment, its sample version is μ_r hat and this is given by summation X_i to the power R , i going from 1 to N divided by N and for r equal to 1, 2 etc, so this is the definition of sample moments. For example; when we put r is equal to 1, we have μ_1 hat that is equal to μ hat that is the sample average and given by 1 by N summation X_i , i going from 1 to N .

Similarly μ_2 hat that is the estimator for mean square value and it is given by 1 by N summation X_i square, i going from 1 to N . So we have defined moments and the sample moments, moments are also known as population moments. Note that E of μ_r hat that is if I take the expectation, because expectation is a linear person we can take insight. So that way this will be 1 by N into summation i going from 1 to N , E of X_i to the power r .

And all are iid, so all will have the same moment therefore we will have N such terms and divided by N and N will get cancelled, we will get E of X_i to the power r . So this means that μ_r hat is an unbiased estimator. Similarly for this unbiased estimator variance of μ_r hat is given by 1 by N square, summation variance of X_i to the power r , because of the independent property, all cross terms will be 0.

So that way we will have simply variance of X_i to the power r divided by N , so that N tends to infinity this quantity will go down to 0, therefore μ_r hat is also consistent. So what we conclude that μ_r hat that is simple moment is an unbiased and a consistent estimator.

(Refer Slide Time: 07:00)

Method of Moments

- ❖ Based on the assumption that the observed data have the sample moments same as the population moments:

$$\hat{\mu}_r = EX_1^r$$

- ❖ MM estimation find k equations relating the first k moments $\mu_1, \mu_2, \dots, \mu_k$ with the parameters $\theta_1, \theta_2, \dots, \theta_k$ and then substitute the moments by the corresponding sample moments $\hat{\mu}_1, \hat{\mu}_2, \dots, \hat{\mu}_k$ in these equations. The solution of the equations give the estimators $\hat{\theta}_1, \hat{\theta}_2, \dots, \hat{\theta}_k$
- ❖ The MM method is also known as the method of substitution

Now we will discuss what is method of moments; it is based on the assumption that the observed data have the sample moments same as the population moments. So whatever sample moments we have that is same as the population moment. So that way μ_r hat is equal to E of X_1 to the power r , this is the example we use in method of moments. Now MM estimation finds k equations relating the first k moments μ_1, μ_2 up to μ_k , with the parameters θ_1, θ_2 up to θ_k .

Then it is substitute the moments, these moments by the corresponding sample moments μ_1 hat, μ_2 hat up to μ_k hat, then we find the solution of the equation the solution of this equation give the estimators θ_1 hat, θ_2 hat up to θ_k hat. So this is the principal part we will find out the first k moments in terms of the parameters, then substitute the moments by the corresponding sample moments and then solve the equations.

The MM method is also known as the method of substitution because you are substituting the true moments by the corresponding sample moments.

(Refer Slide Time: 08:32)

Steps in MM estimation

❖ The following are the steps:

- (1) Express EX_1^r , $r=1,2,\dots,k$ as functions of $\theta_1, \theta_2, \dots, \theta_k$ to get the equations

$$EX_1 = h_1(\theta_1, \theta_2, \dots, \theta_k)$$

$$EX_1^2 = h_2(\theta_1, \theta_2, \dots, \theta_k)$$

$$\vdots$$

$$EX_1^k = h_k(\theta_1, \theta_2, \dots, \theta_k)$$

- (2) Substitute EX_1^r , $r=1,2,\dots,k$ by the corresponding sample moments $\hat{\mu}_r$ to get the modified set of equations

$$\hat{\mu}_1 = h_1(\theta_1, \theta_2, \dots, \theta_k)$$

$$\hat{\mu}_2 = h_2(\theta_1, \theta_2, \dots, \theta_k)$$

$$\vdots$$

$$\hat{\mu}_k = h_k(\theta_1, \theta_2, \dots, \theta_k)$$

Solve the modified set of equations to get the MM estimators $\hat{\theta}_1, \hat{\theta}_2, \dots, \hat{\theta}_k$

We will describe the method in steps, the following are the steps; first you express, the E of X_1 to the power r , r is equal to 1 to up to k as the function of θ_1, θ_2 up to θ_k to get the equation. So these are the equation E of X_1 is equal to some function of θ_1, θ_2 up to θ_k . So that way we write E of X_1 is equal to $h_1 \theta_1, \theta_2$ up to θ_k . Similarly E of X_1 square is equal to $h_2 \theta_1, \theta_2$ up to θ_k and so on up to E of X_1 to the power k that is a function of θ_1, θ_2 up to θ_k and this function will call as h_k .

So that way we have k equations relating the moments with the parameters. Now substitute E of X_1 to the power r , that the r -th moment by the corresponding r -th sample moment to get the modified set of equation. Now we will substitute E of X_1 by $\hat{\mu}_1$, E of X_1 square by $\hat{\mu}_2$ and so on up to knee up X_1^k we will substitute by $\hat{\mu}_k$. Now again we have k equation but this time it is in terms of the sample moments.

Solve the modified set of equation to get the MM estimators $\hat{\theta}_1, \hat{\theta}_2$ up to $\hat{\theta}_k$. So this set of equation there are k unknowns, this k unknowns are the $\hat{\theta}_1, \hat{\theta}_2$ up to $\hat{\theta}_k$. So we can solve this set of equations to get the MM estimators.

(Refer Slide Time: 10:28)

Example: Let $X_1, X_2, \dots, X_N$ are iid with $X_i \sim N(\mu, \sigma^2)$. Find MM estimators for μ and σ^2 .

❖ We have the first two moments,

$$EX_1 = \mu$$

$$EX_1^2 = \sigma^2 + \mu^2$$

❖ Substituting the moments by the sample moments,

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

$$\frac{1}{N} \sum_{i=1}^N x_i^2 = \sigma^2 + \mu^2$$

❖ Solving we get,

$$\hat{\mu}_{MM} = \frac{1}{N} \sum_{i=1}^N x_i \quad \text{and} \quad \hat{\sigma}_{MM}^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \hat{\mu}_{MM})^2$$

We shall consider one example let X_1, X_2 up to X_N are iid independent and identically distributed with each X_i is normally distributed with mean μ and variance σ^2 . Find MM estimators for μ and σ^2 . We have the first two moments for normal distribution, E of X_1 is equal to μ and E of X_1 square is equal to $\sigma^2 + \mu^2$. Now we substitute the moments E of X_1 and E of X_1 square by the corresponding sample moments.

Therefore what we will get $\frac{1}{N} \sum_{i=1}^N x_i$, i going from 1 to N that will be equal to μ , $\frac{1}{N} \sum_{i=1}^N x_i^2$, i going from 1 to N that is equal to $\sigma^2 + \mu^2$. So if we solve this 2 equation in terms of μ and σ^2 we will get, $\hat{\mu}_{MM}$ is equal to $\frac{1}{N} \sum_{i=1}^N x_i$, i going from 1 to N and $\hat{\sigma}_{MM}^2$ that will be equal to $\frac{1}{N} \sum_{i=1}^N (x_i - \hat{\mu}_{MM})^2$, i going from 1 to N .

So that way we got the estimators for mean that is μ and the σ^2 . So this is the estimator for mean, this is the estimator for σ^2 .

(Refer Slide Time: 12:00)

Properties of MM estimators

- ❖ The properties of the MM estimators depend on the nature of $f(x_1, x_2, \dots, x_N; \theta_1, \theta_2, \dots, \theta_K)$.
- ❖ As $\hat{\mu}_r$ is an unbiased and consistent estimator of EX_1^r , the MM estimators will be unbiased and consistent if $\hat{\theta}_r$ is a linear combination of sample moments.
- ❖ Properties of MM estimators are empirically studied through simulations.

We have to look for better estimation rules!

Let us tell briefly about the properties of MM estimators; the properties of MM estimators depend on the measure of this distribution. So for each distribution we will have different properties. As $\hat{\mu}_r$ is an unbiased and a consistent estimator of E of X_1 to the power r that we have already established the MM estimators will be unbiased and consistent if $\hat{\theta}_r$ is a linear combination of sample moments, this is an observation

Usually properties of MM estimators are empirically studied through simulation because the general properties are not obvious we apply numerical simulation to find out this suppose for example bias variance etc and then establish whether it is a consistent estimator, whether it is an unbiased estimator etc. So MM estimator is a very simple estimator it was introduced by Pearson as Bacchus 1901 and it is widely used also but it may not have the good properties of an estimator. So we have to look for a better estimator

(Refer Slide Time: 13:28)

Maximum Likelihood Estimator (MLE)

- ❖ Suppose $X_1, X_2, \dots, X_N$ are random samples with the joint probability density function $f(x_1, x_2, \dots, x_N; \theta) = f(\mathbf{x}; \theta)$ which depends on an unknown non-random parameter θ .
- ❖ Note that $f(\mathbf{x}; \theta)$ is called the **likelihood function**. If $X_1, X_2, \dots, X_N$ are discrete, then the likelihood function will be a joint probability mass function.
- ❖ $L(\mathbf{x}; \theta) = \ln f(\mathbf{x}; \theta)$ is the **log likelihood function**.
- ❖ In discrete case, $f(\mathbf{x}; \theta)$ is replaced by $p(\mathbf{x}; \theta)$.
- ❖ As a function of the random variables, the likelihood and log-likelihood functions are random variables.

Such a rule is the maximum likelihood estimator; suppose X_1, X_2 up to X_N are random samples with the joint probability density function f of x_1, x_2 up to x_N as a function of θ and this we write as f of $\mathbf{x}$ vector as a function of θ . So this is the probability density function and it depends on an unknown non random parameter θ . Suppose we are considering only in the case of simple parameter case, then it depends on an unknown non-random parameter θ .

And note that this $f(\mathbf{x}, \theta)$ is called a likelihood function. If x_1, x_2 up to x_N are discrete, then the likelihood function will be a joint probability mass function that also we know. Now L of $\mathbf{x}, \theta$ that is equal to $\log$ of $f(\mathbf{x}, \theta)$ is the log likelihood function. In discrete case, $f(\mathbf{x}, \theta)$ is replaced by $p(\mathbf{x}, \theta)$ in this expression to find out this a log likelihood function. Now at $f(\mathbf{x}, \theta)$ and $p(\mathbf{x}, \theta)$ they are function of $\mathbf{x}$ and when this $\mathbf{x}$ are varying.

We can consider this to be a function of random variable similarly this to be a function of random variables. Therefore the likelihood and log likelihood functions are also random variables, when we consider for varying $\mathbf{x}$.

(Refer Slide Time: 15:06)

Maximum Likelihood Principle

- ❖ Select that value of θ which maximizes $f(x_1, x_2, \dots, x_N; \theta)$. Thus 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 $\left. \frac{\partial f(x; \theta)}{\partial \theta} \right|_{\hat{\theta}_{MLE}} = 0$

- ❖ Since $L(x; \theta) = \ln(f(x; \theta))$ is a monotonic function of the argument, it is convenient to express the MLE conditions in terms of the log-likelihood function

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

Now we will state the maximum likelihood principle according to this principle select that value of theta which maximizes the likelihood function. Thus maximum likelihood estimator will denoted by theta hat MLE is such an estimator that likelihood at theta hat MLE is greater than equal to the likelihood xN, theta, this is true for all theta. Therefore this is the maximum likelihood principle that likelihood of theta hat MLE is greater than equal to the likelihood of any other theta.

If the likelihood function is differentiable with respect to theta and then theta hat MLE is given by this relation that is partial derivative of likelihood function with respect to theta at theta hat MLE is equal to 0 and now log function is a monotonic function of the argument, therefore this function is also monotonic of this f x, theta. Therefore it is convenient to express the MLE condition in terms of log likelihood function as this.

That is partial derivative of Lx, theta with respect to theta at theta hat MLE is equal to 0. So either we can use this condition or this condition to find out theta hat MLE.

(Refer Slide Time: 16:38)

Maximum Likelihood Principle ...

❖ If we have k unknown parameters given by

$$\theta = [\theta_1 \ \theta_2 \ \dots \ \theta_k]'$$

then the MLE is given by a set of equations.

$$\left. \frac{\partial f(\mathbf{x}; \theta)}{\partial \theta_1} \right|_{\theta_1 = \hat{\theta}_{1,MLE}} = \left. \frac{\partial f(\mathbf{x}; \theta)}{\partial \theta_2} \right|_{\theta_2 = \hat{\theta}_{2,MLE}} = \dots = \left. \frac{\partial f(\mathbf{x}; \theta)}{\partial \theta_k} \right|_{\theta_k = \hat{\theta}_{k,MLE}} = 0$$

❖ In terms of $L(\mathbf{x}; \theta)$ the set of equations are given by

$$\left. \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta_1} \right|_{\theta_1 = \hat{\theta}_{1,MLE}} = \left. \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta_2} \right|_{\theta_2 = \hat{\theta}_{2,MLE}} = \dots = \left. \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta_k} \right|_{\theta_k = \hat{\theta}_{k,MLE}} = 0$$

If we have k unknown parameters given by, theta vector that is equal to theta1, theta 2 up to theta k transpose, this is the parameter vector now. Then the MLE is given by a set of equations and that is del f del theta 1 at theta 1 hat MLE that is equal to 0, similarly del f del theta 2 f theta 2 hat MLE is equal to 0 like that they left del theta k at theta k hat MLE is equal to 0. So these are the conditions for theta hat MLE.

Now in terms of the log-likelihood function also we get a set of equations, so instead of f , now we can use L , del L del theta1 at theta1 hat MLE is equal to 0, del L del theta2 at theta 2 hat MLE is equal to 0 like that up to del L del theta k at theta k hat MLE is equal to 0. So these two conditions if we solve then we will find out the theta hat MLE.

(Refer Slide Time: 17:58)

Example: MLE for the λ parameter of iid Poisson samples

Solution: $X_1, X_2, \dots, X_N$ are iid random variables with $X_i \sim \text{Poi}(\lambda)$. Find MLE for λ .

Solution:

$$p(\mathbf{x}; \lambda) = p(x_1, x_2, \dots, x_N; \lambda)$$

$$= \prod_{i=1}^N \frac{e^{-\lambda} \lambda^{x_i}}{x_i!}$$

$$\therefore L(\mathbf{x}; \lambda) = \ln p(\mathbf{x}; \lambda)$$

$$= -N\lambda + \sum_{i=1}^N x_i \ln \lambda + \text{terms not involving } \lambda$$

$$\left. \frac{\partial L}{\partial \lambda} \right|_{\lambda = \hat{\lambda}_{MLE}} = 0$$

$$\Rightarrow -N + \sum_{i=1}^N \frac{x_i}{\hat{\lambda}_{MLE}} = 0$$

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

We will consider some examples first example is MLE for the λ parameter of the iid Poisson samples. Suppose, X_1, X_2 up to X_N are iid random variables with each X_i is distributed as Poisson with λ parameter. Find MLE for λ . So let us solve this likelihood function is $P(x_1, x_2 \text{ up to } x_N)$ as a function of λ , now each one is distributed as Poisson.

Therefore each distribution is given by e to the power $-\lambda$, λ to the power x_i divided by factorial x_i ; this is the distribution for x_i . So that way we have N distribution because of independent and identically distributed we will get this as product of i going from 1 to N , e to the power of $-\lambda$, λ to the power x_i divided by factorial x_i . So this is the likelihood of x as a function of λ .

Now if we take the logarithm because that exponential terms are there taking logarithm will be easier, so L of x, λ will be $\log$ of p of x, λ and because entrants are there that we will get as $-N\lambda + \sum_{i=1}^N x_i \log \lambda + \text{terms not involving } \lambda$, because we want to do the partial derivative with respect to λ .

So we need not consider the terms which do not involve λ . So if we take the partial derivative with respect to λ at $\lambda \text{ hat MLE}$, we will get 0. So we will take the partial derivative of this will be equal to $-N$ and similarly this part will give you 1 by $\lambda \text{ hat MLE}$. So that way we will get $-N + \sum_{i=1}^N x_i$ divided by partial derivative of this is 1 by λ , so that way it will be $\lambda \text{ hat MLE}$ that will be equal to 0.

So if we solve this we will get that $\lambda \text{ hat MLE}$ is equal to 1 by N summation x_i , i going from 1 to N . So that way we can find out the MLE maximum likelihood estimator for the λ parameter. So this is the sample mean only.

(Refer Slide Time: 20:40)

Example: MLE for multiple parameters

Let $X_1, X_2, \dots, X_N$ be iid random variables with $X_i \sim N(\mu, \sigma^2)$. Find MLE for μ and σ^2 .

Solution:

$$\begin{aligned} f(\mathbf{x}; \mu, \sigma^2) &= f(x_1, x_2, \dots, x_N; \mu, \sigma^2) \\ &= \prod_{i=1}^N \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2}\left(\frac{x_i - \mu}{\sigma}\right)^2} \\ \therefore L(\mathbf{x}; \mu, \sigma^2) &= \ln f(\mathbf{x}; \mu, \sigma^2) \\ &= -N(\ln \sqrt{2\pi}\sigma) - \frac{1}{2} \sum_{i=1}^N \left(\frac{x_i - \mu}{\sigma} \right)^2 \\ \left. \frac{\partial L}{\partial \mu} \right|_{\hat{\mu}_{MLE}} &= 0 \\ \Rightarrow \sum_{i=1}^N (x_i - \hat{\mu}_{MLE}) &= 0 \end{aligned}$$

Second example will consider MLE for multiple parameters; let X_1, X_2 up to X_N are iid random variables with this X_i distributed as a normal distribution with mean μ and variance σ^2 . Find MLE for μ and σ^2 . So we have this likelihood function that is f of x_1, x_2 up to x_N as a function of μ, σ^2 . Now we have iid independent and identically distributed, so we will have a product of N PDFs.

So product i going from 1 to N , $1/\sqrt{2\pi\sigma}$ into e to the power $-\frac{1}{2}\left(\frac{x_i - \mu}{\sigma}\right)^2$ divided by σ^2 . So this is the joint PDF of x_1, x_2 up to x_N as a function of μ and σ^2 , again here exponentially they are so taking logarithm will be beneficial so log likelihood function will be given by $-N \ln \sqrt{2\pi\sigma} - \frac{1}{2}$ of summation i going from 1 to N of $\left(\frac{x_i - \mu}{\sigma}\right)^2$.

Now this is the expression and this we will take the partial derivative with respect to μ and σ^2 . So if I take the partial derivative with respect to μ at $\hat{\mu}_{MLE}$, we will get because these 2 and 2 will get cancelled summation $x_i - \hat{\mu}_{MLE}$, i going from 1 to N , that must be equal to 0, this is one equation.

(Refer Slide Time: 22:32)

Example ...

$$\left. \frac{\partial L}{\partial \sigma^2} \right|_{\hat{\sigma}_{MLE}^2} = 0$$

$$\Rightarrow -\frac{N}{\hat{\sigma}_{MLE}^2} + \frac{\sum_{i=1}^N (x_i - \hat{\mu}_{MLE})^2}{\hat{\sigma}_{MLE}^4} = 0$$

Solving we get

$$\hat{\mu}_{MLE} = \frac{1}{N} \sum_{i=1}^N x_i \quad \text{and}$$

$$\hat{\sigma}_{MLE}^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \hat{\mu}_{MLE})^2$$

Similarly del L del Sigma square at Sigma hat square MLE that is also equal to 0, so from depth we will get - N by Sigma hat MLE square + summation x i - Mu hat MLE whole square, i going from 1 to N divided by Sigma hat to the power 4 MLE is equal to 0. So this is the equation given, so if I take the derivative with respect to Sigma square we will get this expression.

If we solve these two equations part equation is this and second equation is this, we will get Mu hat MLE is equal to 1 by N summation xi, i going from 1 to N. This is the sample mean and similarly this is the sample variance and that is Sigma hat MLE squared is equal to 1 by N summation xi - Mu hat MLE whole square, i going from 1 to N and so that way we get the estimators for mean and variance.

(Refer Slide Time: 23:49)

Example : Non-differentiable likelihood function

Let $X_1, X_2, \dots, X_N$ be iid random samples with the PDF

$$f(x; \theta) = \frac{1}{2} e^{-|x - \theta|} \quad -\infty < x < \infty.$$

Show that *median* $(x_1, x_2, \dots, x_N)$ is the MLE for θ .



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

$$L(\mathbf{x}; \theta) = \ln f(\mathbf{x}; \theta)$$

$$= -N \ln 2 - \sum_{i=1}^N |x_i - \theta|$$

$$\sum_{i=1}^N |x_i - \theta| \text{ is minimized by } \text{median}(x_1, x_2, \dots, x_N)$$

$$\therefore \hat{\theta}_{MLE} = \text{median}(x_1, x_2, \dots, x_N)$$

In the third example, we will consider a non differentiable likelihood function. Let X_1, X_2 up to X_N be iid random variables with the PDF given by this $f(x, \theta)$, θ is equal to $1/2 e^{-|x - \theta|}$, so this is symmetric about θ and this distribution is known as the Laplace distribution. So this is the PDF, so it will be distributed like this compared to Gaussian it will have a longer tail and its mean value is that θ .

Now for this distribution will show that median of x_1, x_2 up to x_N is the MLE for θ , this PDF is given by this because it is a iid independent and identically distributed random variable. So if we have to take the joint PDF that is the product, so this product is given by this. So log likelihood function is given by $-N \log \int_{-\infty}^{\infty} \prod_{i=1}^N f(x_i - \theta) d\theta$, i doing from 1 to N.

Because it is a negative sign is here, we have to minimize and this minimization is done by the median of x_1, x_2 up to x_N , this absolute sum will be minimized if we take θ to be the median. Therefore $\hat{\theta}_{MLE}$ is the median of x_1, x_2 up to x_N . We see that for the Laplace distribution, the MLE for θ is not the mean but the median of x_1, x_2 up to x_n , so that way we considered the how to find out MLE for a non differentiable likelihood function.

(Refer Slide Time: 25:47)

Summary

- ❖ 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

$$\begin{aligned}\hat{\mu}_1 &= h_1(\theta_1, \theta_2, \dots, \theta_k) \\ \hat{\mu}_2 &= h_2(\theta_1, \theta_2, \dots, \theta_k) \\ &\vdots \\ \hat{\mu}_k &= h_k(\theta_1, \theta_2, \dots, \theta_k)\end{aligned}$$
- ❖ MM estimators may or may not satisfy the desired properties of a good estimator.

Let us summarize the lecture today, 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. Method of moment estimation involves relating the moments with the parameters and then substituting the moments by the corresponding sample moments.

Therefore MM estimators are obtained by solving the following set of equations that is $\hat{\mu}_1$ is equal to $\sum_{i=1}^n x_i$, $\hat{\mu}_2$ is equal to $\sum_{i=1}^n x_i^2$ up to $\hat{\mu}_k$ is equal to $\sum_{i=1}^n x_i^k$. So this set of equations if we solve we will get the MM estimators. MM estimators may or may not satisfy the desired properties of a good estimator. To numerical simulation, we can study the properties of MM estimators.

(Refer Slide Time: 27:04)

Summary...

- ❖ 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

$$\left. \frac{\partial f(\mathbf{x}; \theta)}{\partial \theta} \right|_{\hat{\theta}_{MLE}} = 0$$
 or equivalently

$$\left. \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} \right|_{\hat{\theta}_{MLE}} = 0$$
- ❖ If we have k unknown parameters $\theta = [\theta_1 \ \theta_2 \ \dots \ \theta_k]'$ then the MLEs are given by a set of equations:

$$\left. \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta_1} \right|_{\hat{\theta}_{MLE}} = \left. \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta_2} \right|_{\hat{\theta}_{MLE}} = \dots = \left. \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta_k} \right|_{\hat{\theta}_{MLE}} = 0$$

Then we discussed maximum likelihood estimator and $\hat{\theta}_{MLE}$ is such an estimator that this likelihood function at $\hat{\theta}_{MLE}$ is greater than equal to the likelihood function at any other θ that is true all θ , this is the likelihood principle. If the likelihood function is differentiable with respect to θ , then $\hat{\theta}_{MLE}$ is given by this relationship partial derivative of f with respect to θ at $\hat{\theta}_{MLE}$ is equal to 0.

Equivalently the terms of log likelihood function $\frac{\partial L}{\partial \theta}$ at $\hat{\theta}_{MLE}$ is equal to 0. So this is the equation we solve to find out the value of $\hat{\theta}_{MLE}$. If we have k unknown parameters θ is equal to θ_1, θ_2 up to θ_k transpose, this is the parameter vector. Then the MLEs are given by a set of equations that is $\frac{\partial L}{\partial \theta_1}$, $\frac{\partial L}{\partial \theta_2}$ up to $\frac{\partial L}{\partial \theta_k}$ must be equal to 0.

$\frac{\partial L}{\partial \theta_2}$, θ_2 is equal to $\hat{\theta}_2$, that must be equal to 0, like that; $\frac{\partial L}{\partial \theta_k}$, θ_k is equal to $\hat{\theta}_k$ that must be equal to 0. In the next lecture we will see the properties of ML estimators, thank you.

