

Lecture 14
Properties of Maximum Likelihood Estimators (MLE)

Hello students welcome to lecture 14, in this lecture we will discuss properties of maximum likelihood estimators.

(Refer Slide Time: 00:45)

Recall that

- ❖ $\hat{\theta}_{MLE}$ is an estimator such 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
$$\left. \frac{\partial f(\mathbf{x}; \theta)}{\partial \theta} \right|_{\hat{\theta}_{MLE}} = 0$$
- equivalently,
$$\left. \frac{\partial \ln(\mathbf{x}; \theta)}{\partial \theta} \right|_{\hat{\theta}_{MLE}} = 0$$
- ❖ The MLE conditions are extended to multiparameter case.

Recall that, theta hat MLE maximum likelihood estimator is an estimator such that density function we are considering as a function of theta hat MLE is greater than equal to the density function as a function of any other theta, so this is the basic definition and if you likelihood function is differentiable with respect to theta. Then theta hat MLE is given by this expression, del f del theta at theta hat MLE is equal to 0.

Equivalently, we can write in terms of the log likelihood function del L del theta, theta hat MLE is equal to 0. These MLE conditions can be extended to multi-parameter case.

(Refer Slide Time: 01:43)

In this lecture, we will establish some important properties of MLE.

We will now discuss some important properties of MLE.

(Refer Slide Time: 01:46)

Properties of MLE

- ❖ MLEs are backed by elegant mathematical theories. They possess desirable properties of good estimators, particularly for large samples
 - ❖ In many situations, MLE turns out to be MVUE
 - ❖ For some distributions, closed-form solutions of likelihood equations may not exist. In such cases, MLEs may be constructed numerically through iterative algorithms and their properties may also be studied.
- Some properties of MLEs are described next.

MLEs are backed by elegant mathematical theories. They possess desirable properties of good estimators, particularly for large samples. In many situations, MLE turns out to be MVUE minimum variance unbiased estimator. For some distributions, closed form solutions of likelihood equations may not exist. In such cases MLEs may be constructed numerically through iterative algorithms and their properties may be studied.

For example, we may determine the bias-variance etcetera. Some properties of MLEs are described next.

(Refer Slide Time: 02:42)

Properties of MLE

(1) MLE may be biased or unbiased.

In the Example of iid Gaussian samples,

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

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

We can show that

$$E\hat{\mu}_{MLE} = \mu \quad \text{and}$$

$$E\hat{\sigma}_{MLE}^2 = \frac{N-1}{N} \sigma^2$$

$\hat{\mu}_{MLE}$ unbiased estimator
 $\hat{\sigma}_{MLE}^2$ is biased

The first property is an MLE maximum likelihood estimator maybe biased or unbiased. In the example of iid Gaussian samples, we establish that the MLE form μ that is $\hat{\mu}_{MLE}$ is equal to summation x_i , i going from 1 to N divided by N and similarly MLE for Sigma square, $\hat{\sigma}_{MLE}^2$ that is equal to $\frac{1}{N}$ summation i going from 1 to N of $x_i - \hat{\mu}_{MLE}$ whole Square.

And now we can show that if I take the expected value of this then right hand side also I have to take the expected value of its x_i and x_i are iid, so it will have the same mean so therefore we can show that expected value of $\hat{\mu}_{MLE}$ is equal to μ , that is the true parameter. So that way $\hat{\mu}_{MLE}$ is an unbiased estimator. Now, if I take the expected value of $\hat{\sigma}_{MLE}^2$.

Then I will get same way I can expand this and take the expected value then I can show that this expected value is equal to $\frac{N-1}{N}$ times Sigma square. So this is not equal to Sigma square, therefore $\hat{\sigma}_{MLE}^2$ is a biased estimator. So $\hat{\mu}_{MLE}$ unbiased estimator and $\hat{\sigma}_{MLE}^2$ is biased. So that way MLE may be biased or unbiased.

(Refer Slide Time: 04:51)

Properties of MLE

(2) If a sufficient statistic $T(\mathbf{x})$ exists for θ , then $\hat{\theta}_{MLE}$ is a function of $T(\mathbf{x})$.

Proof:

By the factorization theorem,

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

$$\therefore L(\mathbf{x}; \theta) = \ln g(\theta, T(\mathbf{x})) + \ln h(\mathbf{x})$$

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

$$\Rightarrow \left. \frac{\partial}{\partial \theta} (\ln g(\theta, T(\mathbf{x})) + \ln h(\mathbf{x})) \right|_{\hat{\theta}_{MLE}} = 0$$

$$\Rightarrow \left. \frac{\partial}{\partial \theta} (\ln g(\theta, T(\mathbf{x}))) \right|_{\hat{\theta}_{MLE}} = 0$$

Therefore, $\hat{\theta}_{MLE}$ is a function of the sufficient statistic $T(\mathbf{x})$.

Next, one important property if a sufficient statistic Tx exists for θ , then $\hat{\theta}_{MLE}$ is a function of Tx . So if we know that, if sufficient statistic Tx is there then $\hat{\theta}_{MLE}$ must be a function of Tx . This proof is also easy now, by factorization theorem, that likelihood function is product of g of θ Tx into h x , this is the factorization theorem. One part that involved this statistic and θ other part simply functions of x .

Therefore if we take the logarithm, log likelihood function will be log of this part g θ Tx + log of h x . Now, since $\frac{\partial L}{\partial \theta}$, $\hat{\theta}_{MLE}$ is equal to 0, we will get $\frac{\partial}{\partial \theta}$ of log of g θ Tx + log of h x $\hat{\theta}_{MLE}$, it must be equal to 0 and this part does not involve θ , therefore we will simply get $\frac{\partial}{\partial \theta}$ of log of g θ Tx at $\hat{\theta}_{MLE}$ must be equal to 0.

Since $\hat{\theta}_{MLE}$ is the solution of this equation. Therefore $\hat{\theta}_{MLE}$ must be a function of sufficient statistic Tx . So this property we established if a sufficient statistic Tx exists for θ , then $\hat{\theta}_{MLE}$ is a function of Tx .

(Refer Slide Time: 06:44)

Properties of MLE

(3) If an efficient estimator exists, the ML estimator is the efficient estimator.

Suppose an efficient estimator $\hat{\theta}$ exists. Then by Cramer Rao theorem,

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

Note that $\hat{\theta}$ is an MVUE

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

$$\begin{aligned}\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} \bigg|_{\hat{\theta}_{MLE}} &= 0 \\ \Rightarrow C(\hat{\theta} - \hat{\theta}_{MLE}) &= 0 \\ \Rightarrow \hat{\theta}_{MLE} &= \hat{\theta}\end{aligned}$$

Next property is if an efficient estimator exists, the ML estimator is the efficient one. So this is an important property because we always look for an efficient estimator and if we know that it exists then ML estimator will be the efficient estimator. Now, suppose an efficient estimator $\hat{\theta}$ exists, we know that efficient estimator satisfies the Cramer Rao lower bound with equality.

And if we apply the Cramer Rao theorem here, then the condition for efficient estimator is that the partial derivative of the log likelihood function $\frac{\partial L}{\partial \theta}$ must be product of $I(\theta)$ and $\hat{\theta} - \theta$. Where $\hat{\theta}$ is the efficient estimator, this is the condition we get from the Cramer Rao theorem and in that case $\hat{\theta}$ will be a MVUE minimum variance unbiased estimator.

And now let us examine this case suppose one efficient estimator $\hat{\theta}$ exists, now that θ is equal to $\hat{\theta}_{MLE}$ $\frac{\partial L}{\partial \theta}$ is equal to 0. So it implies that this is a constant C into $\hat{\theta} - \hat{\theta}_{MLE}$ must be equal to 0 and this implies that $\hat{\theta}_{MLE}$ is equal to $\hat{\theta}$. So we have established that, if an efficient estimator $\hat{\theta}$ exists then $\hat{\theta}$ must be equal to $\hat{\theta}_{MLE}$.

So this is an important property but here also we presume that the efficient estimator $\hat{\theta}$ exists.

(Refer Slide Time: 08:44)

(5) Invariance Properties of MLE

It is a remarkable property of the MLE and not shared by other estimators. If $\hat{\theta}_{MLE}$ is the MLE of θ and $h(\theta)$ is a function, then $h(\hat{\theta}_{MLE})$ is the MLE of $h(\theta)$.

Proof- We prove the result when $h(\theta)$ is one-to-one.

Suppose $u = h(\theta)$. Then $\theta = h^{-1}(u)$ is given by

$$\frac{\partial f(\mathbf{x}; \theta)}{\partial h(\theta)} = \frac{\partial f(\mathbf{x}; \theta)}{\partial \theta} \times \frac{\partial \theta}{\partial h(\theta)}$$

At $\hat{\theta}_{MLE}$, $\frac{\partial f(\mathbf{x}; \theta)}{\partial \theta} = 0$. Therefore,

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

$h(\hat{\theta}_{MLE})$ is the MLE of $h(\theta)$.

Next property is the invariance properties of MLE. It is a remarkable property of the MLE and not shared by other estimators, so if $\hat{\theta}_{MLE}$ is the MLE of θ and $h(\theta)$ is a function, then $h(\hat{\theta}_{MLE})$ is the MLE of $h(\theta)$. So here we know suppose $\hat{\theta}_{MLE}$ is the MLE of θ . Now, if θ is another function and for this is the maximum likelihood estimator will be a h of $\hat{\theta}_{MLE}$, so that is the invariance.

The under transformation that maximum likelihood point is invariant so whatever $\hat{\theta}_{MLE}$ is there at that point corresponding function is $\hat{\theta}_{MLE}$ will be the MLE for $h(\theta)$. We proof to result when $h(\theta)$ is 1 to 1 this is a simple case, suppose u is equal to $h(\theta)$ then θ is equal to $h^{-1}(u)$, because 1 to 1 it is invertible. Now $\frac{\partial f}{\partial h(\theta)}$ that will be equal to $\frac{\partial f}{\partial \theta}$ multiplied by $\frac{\partial \theta}{\partial h(\theta)}$, now this quantity exists.

And this is a invertible function, so we can find out this $\hat{\theta}_{MLE}$ this quantity $\frac{\partial f}{\partial h(\theta)}$ will be equal to 0 and this is usually nonzero. Therefore we will get $\frac{\partial f}{\partial \theta}$, $\frac{\partial \theta}{\partial h(\theta)}$ at $\hat{\theta}_{MLE}$ is equal to 0. So this will imply that this maximum likelihood estimator for $h(\theta)$ will be at $h(\hat{\theta}_{MLE})$, so $h(\hat{\theta}_{MLE})$ is the MLE of $h(\theta)$.
(Refer Slide Time: 10:57)

(5) Invariance Properties of MLE...

❖ We have proved the invariance property for the simple case when $h(\theta)$ is one to one and differentiable. However, the result is true also when $h(\theta)$ is many-to-one.

❖ Invariance to a transformation is a remarkable property, not shared by other estimators.

Example- In our example of iid Gaussian samples,

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

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

We have proved the invariance property for the simple case when $h(\theta)$ is one-to-one and differentiable. However the result is true also when $h(\theta)$ is many-to-one. Invariance to a transformation is a remarkable property, not shared by any other estimators. Example in our example of iid Gaussian sample, suppose $\hat{\sigma}_{MLE}^2$ that is the MLE for variance is equal to $\frac{1}{N} \sum_{i=1}^N (x_i - \hat{\mu}_{MLE})^2$.

Therefore what will be the $\hat{\sigma}_{MLE}$, so that is the MLE for standard deviation that is $\hat{\sigma}_{MLE}$ will be given by square root of this MLE expression, therefore this will be the square root of $\frac{1}{N} \sum_{i=1}^N (x_i - \hat{\mu}_{MLE})^2$, i going from 1 to N . So this is the MLE for standard deviation. Once we know the MLE for variance we can find out the MLE for standard deviation using the invariance property.

(Refer Slide Time: 12:15)

Example:

Suppose $X_1, X_2, \dots, X_N$ are iid $B(1, \theta)$ random variables. Find the MLE for θ and hence the MLE for $\text{var}(X_1)$

$$p(\mathbf{x}; \theta) = \theta^{\sum_{i=1}^N x_i} (1-\theta)^{N - \sum_{i=1}^N x_i}$$

$$\therefore L(\mathbf{x}; \theta) = \sum_{i=1}^N x_i \ln \theta + \left(N - \sum_{i=1}^N x_i \right) \ln(1-\theta)$$

Applying the MLE condition, we get

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

Now, $\text{var}(X_1) = \theta(1-\theta)$

$$\therefore \text{var}(X_1)_{MLE} = \hat{\theta}_{MLE}(1 - \hat{\theta}_{MLE})$$

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

We will consider another example suppose X_i 's are iid we wanted this is the symbol for Bernoulli random variable with success parameter θ , so X_i 's are $B, 1, \theta$ random variables. Find the MLE for θ and hence the MLE for variance of X_1 . So here X_i 's are iid Bernoulli. Now p of x θ , x is the vector x_1, x_2 up to x_N . So P of x θ will be θ to the power the number of success is summation x_i, i going from 1 to N , because x_i takes value either 1 for success or 0 for failure.

So that way summation x_i will be the number of success. so θ to the power of summation x_i, i going from 1 to N into $1 - \theta$ to the power $N - \text{summation } x_i, i$ going from 1 to N , this is the likelihood function. So if we take the logarithm then we will get the log likelihood function and this will be given by summation x_i, i going from 1 to N times $\log$ of θ + $N - \text{summation } x_i, i$ going from 1 to N times $\log$ of $1 - \theta$.

So now we can take the partial derivative with respect to θ and make it equal to 0 and we will get $\hat{\theta}$ MLE is equal to $1/N$ into summation x_i, i going from 1 to N . So this is the $\hat{\theta}$ MLE for this parameter θ in the Bernoulli random variable. Now variance of X_1 is equal to $\theta(1 - \theta)$, which is a function of θ . You can consider any variable X_1, X_2 or any X_i because they are iid.

Therefore we can find out the MLE for variance of X_1 , that is variance of X_1 MLE, that is equal to we will substitute $\hat{\theta}$ MLE, so $\hat{\theta}$ MLE into $1 - \hat{\theta}$ MLE were $\hat{\theta}$ MLE is given by the likelihood estimator of θ given by this expression. So that way we can change the application of invariance property of maximum likelihood estimator and this have many uses in signal processing.

(Refer Slide Time: 14:56)

Large sample properties of MLE

- ❖ We saw that if an efficient estimator exists, the MLE is the efficient estimator. Another attractive feature of the MLE is that its behavior becomes better and better with more number of samples.
- ❖ The asymptotic properties of MLE holds under the regularity conditions like those applied in deriving the Cramer Rao theorem. These conditions are usually satisfied in practice
- ❖ The MLE is asymptotically unbiased and efficient. Thus, for large N , the MLE is approximately efficient.

We will explain this property as follows;

Now we will discuss large sample properties of MLE; we saw that if an efficient estimator exists, the MLE is the efficient estimator. Another attractive feature of MLE is that its behavior becomes better and better with more number of samples. So that will establish now. The asymptotic properties of MLE holds under the regularity conditions like those applied in deriving Cramer Rao theorem, we will not discuss about those regularity conditions but these conditions are usually satisfied in practice.

The MLE is asymptotically unbiased and efficient this is a remarkable property, thus for large N and MLE is approximately efficient. asymptotic property means as N tends to infinity therefore when we have large N , the MLE is approximately efficient. We will explain the asymptotic efficiency property.

(Refer Slide Time: 16:03)

Asymptotic efficiency MLE

- ❖ To show that for large N ,

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

We have,

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

$$\begin{aligned} \therefore 0 &= \frac{\partial}{\partial \theta} L(\mathbf{x}; \hat{\theta}_{MLE}) \\ &= \frac{\partial}{\partial \theta} L(\mathbf{x}; \theta + \hat{\theta}_{MLE} - \theta) \\ &= \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} + (\hat{\theta}_{MLE} - \theta) \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta^2}, \quad \theta < \theta_1 < (\hat{\theta}_{MLE}) \end{aligned}$$

where we have applied the mean-value theorem.

$$\therefore \frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} = - \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta^2} (\hat{\theta}_{MLE} - \theta)$$

We have to show that for large N , $\frac{\partial L}{\partial \theta}$ of $L(\mathbf{x}, \theta)$ is equal to $I(\theta)(\hat{\theta}_{MLE} - \theta)$. This is the one we want to solve. If we can show this, then this $\hat{\theta}_{MLE}$ will be asymptotically efficient. Now we will start with the MLE condition $\frac{\partial L}{\partial \theta}$ at $\hat{\theta}_{MLE}$ is equal to 0, therefore 0 will be equal to $\frac{\partial^2 L}{\partial \theta^2}$ of $L(\mathbf{x}, \hat{\theta}_{MLE})$. Now at $\hat{\theta}_{MLE}$, we can write like this $\theta + \hat{\theta}_{MLE} - \theta$ where θ is the original value of the parameter.

So we will write this $\hat{\theta}_{MLE}$ by this manipulation $\theta + \hat{\theta}_{MLE} - \theta$. Now we can apply the mean value theorem to this expression, so we can write this as $\frac{\partial L}{\partial \theta}$ of $\theta + \hat{\theta}_{MLE} - \theta$ into because we are considering this function so derivative of this function will be $\frac{\partial^2 L}{\partial \theta^2}$ at a point θ_1 where θ_1 lies between θ , θ_1 , $\hat{\theta}_{MLE}$.

So this expression we get by applying the mean value theorem, so this quantity $\frac{\partial L}{\partial \theta}$ of $\mathbf{x}, \hat{\theta}_{MLE}$ is same as $\frac{\partial^2 L}{\partial \theta^2}$ of $\mathbf{x}, \theta + \hat{\theta}_{MLE} - \theta$ into $\hat{\theta}_{MLE} - \theta$ squared upon θ_1 , for some θ_1 line between θ , θ_1 , $\hat{\theta}_{MLE}$. Now since left-hand side is 0, so we can write $\frac{\partial L}{\partial \theta}$ is equal to minus, if we take it to the other side - $\frac{\partial^2 L}{\partial \theta^2}$ into $\hat{\theta}_{MLE} - \theta$. So we got one expression like this.

(Refer Slide Time: 18:03)

Asymptotic efficiency MLE....

We have

$$\frac{\partial L(\mathbf{x}; \theta)}{\partial \theta} = -\frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta^2} (\hat{\theta}_{MLE} - \theta)$$

Under the iid assumption, we can apply the weak law of large numbers to show that

$$\frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta^2} \text{ converges in probability to } E \frac{\partial^2 L(\mathbf{x}; \theta)}{\partial \theta^2} = -I(\theta).$$

Therefore, we can write,

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

$$\lim_{N \rightarrow \infty} \text{var}(\hat{\theta}_{MLE}) = \frac{1}{I(\theta)}$$

Thus $\hat{\theta}_{MLE}$ is asymptotically efficient. The same relation also explains the asymptotic unbiasedness of $\hat{\theta}_{MLE}$.

We have established that, $\frac{\partial L}{\partial \theta}$ is equal to - $\frac{\partial^2 L}{\partial \theta^2}$ times $\hat{\theta}_{MLE} - \theta$. Now under the iid assumption because we are assuming that X_i 's are iid, we can apply the weak law of large numbers to this quantity, actually this quantity can be expanded

at the sum of individual log likelihood functions. So that way it is a sum of some large number of log likelihood functions.

Therefore we can apply the weak law of large numbers which means that this quantity $\frac{1}{N} \sum_{i=1}^N \log \ell(\theta_i)$ will converge in probability to the corresponding expected value of $\log \ell(\theta)$ and which is equal to $-\frac{1}{2} I(\theta)$. Therefore we can write, $\frac{1}{N} \sum_{i=1}^N \log \ell(\theta_i) = -\frac{1}{2} I(\theta)$. So we have established that for large N , when N tends to infinity $\frac{1}{N} \sum_{i=1}^N \log \ell(\theta_i)$ can be expressed as $-\frac{1}{2} I(\theta)$ into $\hat{\theta}_{MLE}$ - θ . So we have established that for large N , when N tends to infinity $\frac{1}{N} \sum_{i=1}^N \log \ell(\theta_i)$ can be expressed as $-\frac{1}{2} I(\theta)$ into $\hat{\theta}_{MLE}$ - θ , which means that $\hat{\theta}_{MLE}$ is the efficient estimator.

And if I consider the Cramer Rao, therefore limit for the variance of $\hat{\theta}_{MLE}$ as N tends to infinity will be given by $\frac{1}{N I(\theta)}$, where $I(\theta)$ is the Fisher information statistic and it is given by this expression $E \left[\left(\frac{\partial \log \ell(\theta)}{\partial \theta} \right)^2 \right] = I(\theta)$. Thus $\hat{\theta}_{MLE}$ is asymptotically efficient. The same relation also explains the asymptotic unbiasedness of $\hat{\theta}_{MLE}$, because you see that when this condition is satisfied in that situation $\hat{\theta}_{MLE}$ is the MVUE, minimum variance unbiased estimator.

So this condition implies that this $\hat{\theta}_{MLE}$ will be also asymptotically unbiased. So we have established that $\hat{\theta}_{MLE}$ is asymptotically efficient and it is asymptotically unbiased.

(Refer Slide Time: 20:25)

Consistency of MLE

Recall that

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

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

- ❖ Further, if $\hat{\theta}$ is unbiased and $\lim_{N \rightarrow \infty} \text{var}(\hat{\theta}) = 0$.

It can be shown that under the regularity conditions, $\hat{\theta}_{MLE}$ is a consistent estimator. We omit the proof.

- ❖ The desired properties of $\hat{\theta}_{MLE}$ under large sample conditions make MLE an attractive estimator for the signal processing communities.
- ❖ These properties can be easily extended to multi-parameter case.

Last property we shall consider the consistency of MLE; recall that an estimator $\hat{\theta}$ is a consistent estimator of θ , if $\hat{\theta}$ converges in probability to θ in other words limit

N tends to infinity of the probability that $\hat{\theta} - \theta$ is greater than equal to ϵ will be always equal to 0 for any ϵ greater than 0. So this is the condition for convergence in probability, therefore with this definition we can take whether $\hat{\theta}_{MLE}$ is consistent.

Further, if $\hat{\theta}$ is an unbiased estimator then the condition for consistency is this limit of variance of $\hat{\theta}$ as N tends to infinity is equal to 0. If this happens then $\hat{\theta}$ will be a consistent estimator. Now it can be shown that under the regularity conditions $\hat{\theta}_{MLE}$ is consistent estimator, we can show that it satisfies this property limit of probability of $\hat{\theta} - \theta$ greater than equal to any ϵ as n tends to infinity will be equal to 0.

But we will omit this proof but what we claim is that $\hat{\theta}_{MLE}$ is a consistent estimator under some regularity conditions but those regularity conditions are satisfied in practice. The desired properties of $\hat{\theta}_{MLE}$ under large sample condition make MLE an attractive estimator for the signal processing communities because we have seen that if we consider general properties of $\hat{\theta}_{MLE}$ then they are not very attractive because we have to presume something.

But here these large sample properties assume that the $\hat{\theta}_{MLE}$ is asymptotically efficient and it is a consistent estimator. So these properties also can be extended to multi-parameter cases.

(Refer Slide Time: 22:44)

Consistency of MLE....

In the Example of iid Gaussian samples,

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

We can show that

$$E\hat{\mu}_{MLE} = \mu \quad \text{and}$$

$$\text{var}\hat{\mu}_{MLE} = \frac{\sigma^2}{N}$$

$$\therefore \lim_{N \rightarrow \infty} \text{var}\hat{\mu}_{MLE} = 0$$

Thus, $\hat{\mu}_{MLE}$ is a consistent estimator.

let us give an example, in the example of iid Gaussian samples; θ $\hat{\mu}$ MLE is given by $\frac{1}{N} \sum_{i=1}^N x_i$, i going from 1 to N , this is the MLE for μ . Now we can show that E of $\hat{\mu}$ MLE is equal to μ , it is a unbiased estimator and variance of $\hat{\mu}$ MLE is equal to σ^2 by and using the iid property. We can establish this therefore limit of $\hat{\mu}$ MLE will be equal to 0.

We simply that $\hat{\mu}$ MLE is a consistent estimator, so we have shown that maximum likelihood estimator for μ is a consistent estimator. This is a general property.

(Refer Slide Time: 23:34)

Summary

- (1) MLE may be biased or unbiased.
- (2) If a sufficient statistic $T(\mathbf{x})$ exists for θ , then $\hat{\theta}_{MLE}$ is a function of $T(\mathbf{x})$.
- (3) If an efficient estimator exists, the MLE estimator is the efficient estimator. Thus, if

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

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

To summarize MLE may be biased or unbiased. If a sufficient statistic Tx exists for θ , then $\hat{\theta}$ MLE is a function of Tx . If an efficient estimator exists, then MLE estimator is the efficient estimator. Thus if we can write that $\frac{\partial L}{\partial \theta}(\mathbf{x}; \theta) = I(\theta)(\hat{\theta} - \theta)$, then this $\hat{\theta}$ must be equal to $\hat{\theta}$ MLE.

(Refer Slide Time: 24:08)

Summary...

(4) The MLE is asymptotically unbiased and efficient. Thus, for large N , the MLE is approximately efficient.

(5) $\hat{\theta}_{MLE}$ is a consistent estimator.

- ❖ The desired properties of $\hat{\theta}_{MLE}$ under large sample conditions make MLE an attractive estimator for the signal processing communities.
- ❖ These properties can be easily extended to multi-parameter case.

Then we establish the large sample properties of MLE, MLE is asymptotically unbiased and efficient. Thus, for a large number of samples MLE satisfy the Cramer Rao lower bound with equality. Thus, for large N , MLE is approximately efficient $\hat{\theta}_{MLE}$ is also a consistent estimate, that we did not proof but this is also a general property of maximum likelihood estimator.

The desired properties of $\hat{\theta}_{MLE}$ under large sample conditions make MLE an attractive estimator for signal processing communities. Many signal processing application we use for example there are applications of mean, variance, median, etcetera, where the MLE estimators are used. These properties of MLE can be extended to multi parameter case. Thank you.