

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

Lecture No 16
Bayesian estimators 2

Hello students, welcome to the lecture on Bayesian estimators 2.

(Refer Slide Time: 00:45)

Let us recall

- ❖ Unlike MLE, Bayesian estimators assume the parameter θ to be a random variable with the prior PDF $f(\theta)$ or a prior PMF $p(\theta)$.
- ❖ The estimator uses the posterior PDF $f(\theta/\mathbf{x})$ using the Bayes rule.
- ❖ We associate a *cost function* or a *loss function* $C(\theta - \hat{\theta})$ with the estimator $\hat{\theta}$.
- ❖ A Bayesian estimator solves the optimization problem

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

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

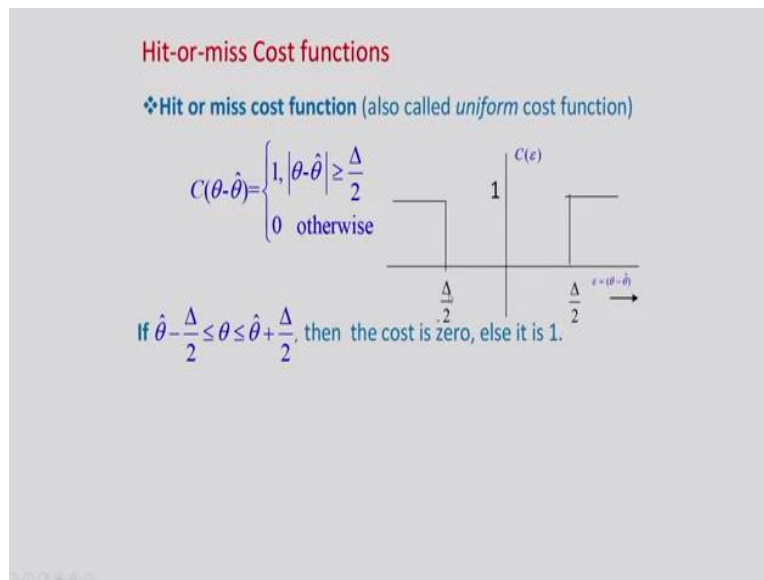
$$\hat{\theta}_{\text{MMSE}} = \int_{-\infty}^{\infty} \theta f(\theta/\mathbf{x}) d\theta = E(\Theta/\mathbf{x})$$

Let us recall unlike MLE Bayesian estimators assume the parameters θ to be a random variable with the prior PDF $f(\theta)$ or prior PMF $p(\theta)$. The estimator uses the posterior PDF $f(\theta/\mathbf{x})$ of θ given $\mathbf{x}$ or posterior PMF $p(\theta/\mathbf{x})$ of θ given $\mathbf{x}$ using the Bayes rule. We associate a cost function or a loss function $C(\theta - \hat{\theta})$ with the estimator $\hat{\theta}$. The cost function is a function of error.

A Bayesian estimator solves the optimization problem, this is the optimization problem minimize $\bar{C}(\hat{\theta})$ over $\hat{\theta}$. $\bar{C}(\hat{\theta})$ is equal to integration from minus infinity to infinity of $C(\theta - \hat{\theta}) f(\mathbf{x}, \theta) d\mathbf{x} d\theta$. This is the average cost function we want to minimize. The MMSE estimator minimum mean square error estimator minimizes the mean square error and is given by this expression $\hat{\theta}_{\text{MMSE}} = \int_{-\infty}^{\infty} \theta f(\theta/\mathbf{x}) d\theta = E(\Theta/\mathbf{x})$. Remember C is equal to integration from minus infinity to infinity of $\theta f(\theta/\mathbf{x}) d\theta$.

That is the conditional expectation update and the variable theta given X so MMSE is given by the conditional expectation of theta given X. In this lecture I will focus on another Bayesian estimator known as the maximum a posteriori probability map estimator. I will also introduce another estimator known as the minimum absolute error MAE estimator.

(Refer Slide Time: 02:40)



Let us start with hit or miss cost functions hit or miss cost function also known as uniform cost function is given by this C of $\theta - \theta_{\text{hat}}$ is equal 1, if the absolute value of $\theta - \theta_{\text{hat}}$ is greater than equal to $\Delta/2$ is 0. Otherwise, so this is the plot of the hit or miss cost function if error is within plus minus $\Delta/2$ then cost is 0, otherwise cost this 1. In other words if parameter θ lies between θ_{hat} minus $\Delta/2$.

And θ_{hat} plus $\Delta/2$ then the cost is 0, else cost is 1 and here this Δ is very small quantity.

(Refer Slide Time: 03:34)

Maximum a posteriori probability (MAP) estimator

Bayesian Risk $\bar{C}(\hat{\theta}) = EC(\theta - \hat{\theta}) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} C(\theta - \hat{\theta}) f(\mathbf{x}, \theta) d\mathbf{x} d\theta$

❖ The Bayesian estimation problem

$$\begin{aligned} \text{Minimize over } \hat{\theta} \quad & \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} C(\theta - \hat{\theta}) f(\theta | \mathbf{x}) f(\mathbf{x}) d\theta d\mathbf{x} \\ & = \int_{-\infty}^{\infty} \left(\int_{-\infty}^{\infty} C(\theta - \hat{\theta}) f(\theta | \mathbf{x}) d\theta \right) f(\mathbf{x}) d\mathbf{x} \end{aligned}$$

❖ As earlier, we have to minimize the inner integral

$$\int_{-\infty}^{\infty} C(\theta - \hat{\theta}) f(\theta | \mathbf{x}) d\theta \quad \text{with respect to } \hat{\theta}.$$

❖ The Bayesian estimation problem for the hit-or-miss cost function is to

$$\text{Minimize over } \hat{\theta} \quad \int_{\hat{\theta} - \frac{\Delta}{2}}^{\hat{\theta} + \frac{\Delta}{2}} f(\theta | \mathbf{x}) d\theta + \int_{\hat{\theta} + \frac{\Delta}{2}}^{\infty} f(\theta | \mathbf{x}) d\theta$$

Now corresponding to this hit-or-miss cost function the Bayesian risk function is given by expected value of $C(\theta - \hat{\theta})$, that is integration from minus infinity to infinity minus infinity to infinity of $C(\theta - \hat{\theta}) f(\theta | \mathbf{x}) f(\mathbf{x}) d\theta d\mathbf{x}$. And now as earlier we will take this $f(\mathbf{x}) d\mathbf{x}$ outside, so we will get the minimization problem at the minimization of this integral.

Integral from minus infinity to infinity integral from minus infinity to infinity $C(\theta - \hat{\theta}) f(\theta | \mathbf{x}) f(\mathbf{x}) d\theta d\mathbf{x}$, so we have one inner integral and one outer integral. Since $f(\mathbf{x})$ is always greater than equal to 0 therefore the double integral will be minimized, whenever the inner integral is minimized. So therefore our minimization problem is to find a minimum of this integral integral from minus infinity to infinity $C(\theta - \hat{\theta}) f(\theta | \mathbf{x}) d\theta$ with respect to $\hat{\theta}$.

Now using this cost function we have the cost function value at this interval is zero therefore we can write the minimizes some problem as minimize over $\hat{\theta}$ integral from $\hat{\theta} - \frac{\Delta}{2}$ to $\hat{\theta} + \frac{\Delta}{2}$ of $f(\theta | \mathbf{x}) d\theta$ plus integral from $\hat{\theta} + \frac{\Delta}{2}$ to infinity of $f(\theta | \mathbf{x}) d\theta$.

(Refer Slide Time: 05:38)

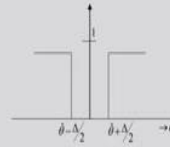
Maximum a posterior (MAP) estimator

The Bayesian estimation problem is now

$$\underset{\hat{\theta}}{\text{Minimize}} \int_{-\infty}^{\hat{\theta} - \frac{\Delta}{2}} f(\theta | \mathbf{x}) d\theta + \int_{\hat{\theta} + \frac{\Delta}{2}}^{\infty} f(\theta | \mathbf{x}) d\theta$$

This is equivalent to minimizing

$$1 - \int_{\hat{\theta} - \frac{\Delta}{2}}^{\hat{\theta} + \frac{\Delta}{2}} f(\theta | \mathbf{x}) d\theta$$



Or maximizing the integral

$$\int_{\hat{\theta} - \frac{\Delta}{2}}^{\hat{\theta} + \frac{\Delta}{2}} f(\theta | \mathbf{x}) d\theta \approx f(\theta | \mathbf{x}) \Big|_{\theta = \hat{\theta}} \cdot \Delta \text{ for small } \Delta$$

The approximation will be large when $f(\theta | \mathbf{x}) \Big|_{\theta = \hat{\theta}}$ is large. As $\Delta \rightarrow 0$, $f(\theta | \mathbf{x}) \Big|_{\theta = \hat{\theta}}$ is maximum when it is equal to the maximum of $f(\theta | \mathbf{x})$.

Thus the Bayesian estimation problem is now minimize over theta hat integral from minus infinity to theta hat minus Delta by 2 F of theta given X D theta plus integral from theta hat plus Delta by 2 to infinity f of theta given X D theta. Now this integral is equivalent to this expression because if we integral suppose this expression from minus infinity to infinity, will get 1, because this is a density function.

Therefore, the this minimization problem can be expressed as the equivalent minimization problem of this integral 1 minus integral theta hat minus Delta by 2 2 theta hat plus Delta by 2 F of theta given X D theta. So now we have to minimize this part since this is a difference, so this minimum will be same as maximizing the integral this integral if we maximize then this difference will be minimum therefore we have to maximize the integral theta hat minus Delta Y to 2 theta hat plus Delta by 2 F of theta given X D theta.

And this we can approximate as F of theta given X at point ax equal to theta hat into Delta for small Delta, because our Delta is very small therefore this integral we can approximate by this expression. Now consider this approximation Delta is a constant therefore this approximation will be large whenever this quantity is large and as delta tends to 0 f of theta given X at theta is equal to theta hat is maximum when it is equal to the maximum of F of theta given X.

(Refer Slide Time: 07:55)

MAP estimator...

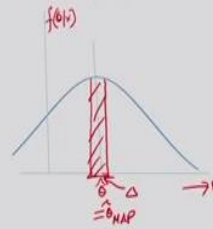
- Thus, the estimator is the value of θ for which $f(\theta | \mathbf{x})$ is maximum and is called the MAP estimator.
- This estimator is denoted by $\hat{\theta}_{MAP}$.

$$\text{Thus, } \hat{\theta}_{MAP} = \arg \max_{\theta} f(\theta | \mathbf{x})$$

Now

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

$$\begin{aligned} \therefore \hat{\theta}_{MAP} &= \arg \max_{\theta} \frac{f(\theta) f(\mathbf{x} | \theta)}{f(\mathbf{x})} \\ &= \arg \max_{\theta} f(\theta) f(\mathbf{x} | \theta) \text{ as } f(\mathbf{x}) \text{ does not involve } \theta \end{aligned}$$



We will illustrate with this diagram, suppose this is my theta and this is my F of theta given X. And now this density function has a maximum at this point now if I put my theta hat here theta hat then this area for this area this interval is Delta so this area will be maximum if I put theta hat at this point corresponding to this maximum and therefore the estimator is the value of theta for which F of theta given X is maximum, and this called the MAP estimator.

Maximum is aposterior probability estimator, we will denote this estimator by theta hat MAP, because this is corresponding to the maximum a-posteriori PDF. So this point is theta hat is equal to theta hat MAP. So this point is equal to theta hat map the theta hat map is argument over theta of max f of theta given X. Now f of theta given X that is equal to f theta into f of X given theta divided by F X.

Therefore theta hat map will be arguing theta of max f theta into f of of X given theta divided by F X. Now this FS does not involve any theta they are prototype map is argmax theta of f theta into f of x given theta.

(Refer Slide Time: 09:49)

MAP equations

❖ If $f(\theta|\mathbf{x})$ is differentiable, then $\hat{\theta}_{MAP}$ is given by

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

❖ Equivalently, $\hat{\theta}_{MAP}$ is given by the solution of

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

The above two equations are known as the MAP equations.

Now if F of θ given X is differentiable, the conditional PDF or a posterior PDF is differentiable then θ at map will be given by $\frac{\partial F}{\partial \theta}$ at θ at map is equal to 0. So partial derivative of this a posterior PDF and $\hat{\theta}_{MAP}$ is equal to 0 or in terms of the log likelihood function we can write $\frac{\partial L}{\partial \theta}$ at $\hat{\theta}_{MAP}$ is equal to 0, either we can solve this equation or this equation. The above two equations are known as the map equations.

(Refer Slide Time: 10:40)

Example

Consider single observation X that depends on a random parameter θ . Suppose

$$f(\theta) = \lambda e^{-\lambda\theta} \text{ for } \theta \geq 0, \lambda > 0 \text{ and } f(x|\theta) = \theta e^{-\theta x} \text{ } x > 0$$

Find the MAP estimation for θ .

$$\text{Solution: We have } f(\theta|x) = \frac{f(\theta)f(x|\theta)}{f(x)}$$

$$\therefore \ln(f(\theta|x)) = \ln(f(\theta)) + \ln(f(x|\theta)) - \ln f(x)$$

Therefore, MAP estimator is given by.

$$\begin{aligned} \left. \frac{\partial}{\partial \theta} \ln f(\theta) + \frac{\partial}{\partial \theta} \ln f(x|\theta) \right|_{\hat{\theta}_{MAP}} &= 0 \\ \Rightarrow \left. \frac{\partial}{\partial \theta} (\ln \lambda - \lambda\theta) + \frac{\partial}{\partial \theta} (\ln \theta - \theta x) \right|_{\hat{\theta}_{MAP}} &= 0 \\ \Rightarrow -\lambda + \frac{1}{\hat{\theta}_{MAP}} - x &= 0 \\ \Rightarrow \hat{\theta}_{MAP} &= \frac{1}{\lambda + X} \end{aligned}$$

Example consider single observation X that depends on a random parameter θ , suppose F θ is equal to λ into e to the power minus $\lambda\theta$ for θ greater than equal to 0 and λ greater than zero and the likelihood $f(x|\theta)$ is equal to θ into e to the power minus θx , $x > 0$. find the map estimator for θ here λ is assumed

to be known we have the a posteriori PDF $f(\theta | x)$ given X that is equal to $F(\theta | x)$ into $f(\theta | x)$ given θ divided by $F(X)$.

And taking the logarithm $\log$ of $F(\theta | x)$ equal to $\log$ of $F(\theta | x)$ plus $\log$ of $F(x | \theta)$ minus $\log$ of $F(X)$ so therefore map estimator is given by as Toller del del θ of $\log$ of $f(\theta | x)$ plus del $L(\theta)$ of $\log$ of x given θ at θ hat map must be equal to 0. So now we have to take the partial derivative or be given to both PDF is given here, so we can find out $\log$ of $f(\theta | x)$. Similarly we can find out $\log$ of this expression.

And if we take the partial derivative we can write del del θ of $\log$ lambda minus minus lambda θ plus del Del θ of $\log$ theta minus theta X theta hat map is equal to 0. So we will take the partial derivative and then put θ is equal to θ hat map we will get minus lambda plus 1 y Tita at map minus X is equal to 0. So from this we can find out θ hat map and the estimator is given by 1 by lambda plus X . So this is this maximum a posteriori estimator.

(Refer Slide Time: 12:55)

Relation between ML and MAP estimators

❖ From $f(\theta | x) = \frac{f(\theta)f(x|\theta)}{f(x)}$

$$\ln(f(\theta | x)) = \ln(f(\theta)) + \ln(f(x|\theta)) - \ln(f(x))$$

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

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

❖ MAP equation can be written as $\left. \frac{\partial \ln(f(\theta))}{\partial \theta} \right|_{\theta_{MAP}} + \left. \frac{\partial \ln(f(x|\theta))}{\partial \theta} \right|_{\theta_{MAP}} = 0$

❖ Compare the above with the ML equation $\left. \frac{\partial \ln(f(x|\theta))}{\partial \theta} \right|_{\theta_{ML}} = 0$

There is additional contribution of the prior PDF.

❖ When $f(\theta)$ is sufficiently flat, $\ln(f(\theta))$ will also be flat so that

$$\left. \frac{\partial \ln(f(\theta))}{\partial \theta} \right|_{\theta_{MAP}} \approx 0. \text{ In that case } \hat{\theta}_{MAP} \approx \hat{\theta}_{MLE}$$

Now let us see the relation between ML and MAP estimators from $f(\theta | x)$ given x equal to $f(\theta | x)$ into $f(x | \theta)$ given θ divided by $f(x)$, we can write $\log$ of $F(\theta | x)$ given X is equal to $\log$ of $F(\theta | x)$ plus $\log$ of $F(x | \theta)$ minus $\log$ of $F(X)$. Now this partial derivative of this with respect to θ is 0 that will imply that $\log$ of $f(\theta | x)$ del θ $\log$ of $f(x | \theta)$ del θ at θ hat map must be equal to 0.

So therefore we have this MAP equation $\frac{\partial}{\partial \theta} \log f(\theta) + \frac{\partial}{\partial \theta} \log f(x|\theta)$ at θ_{MAP} must be equal to 0. And if we examine the ML equation this is simply $\frac{\partial}{\partial \theta} \log \text{likelihood function } f(x|\theta)$ MLE is equal to 0, so we see that this expression contains this term but it has one additional term. There is additional contribution from the prior PDF.

This is the contribution from the prior PDF when $f(\theta)$ is sufficiently flat in that case $\log f(\theta)$ will be also flat so that this quantity $\frac{\partial}{\partial \theta} \log f(\theta)$ will be approximately equal to 0. So this will be approximately 0, so that θ_{MAP} is approximately will be equal to θ_{MLE} . So when $f(\theta)$ is sufficiently flat that means its slope is very small in that case $\log f(\theta)$ will be also flat.

So that its slope will be also very small nearly zero in that case θ_{MAP} will be equal to θ_{MLE} . That means if the parameter has more uncertainty in that case θ_{MAP} will be closer to θ_{MLE} .

(Refer Slide Time: 15:17)

Example
 Let $X_1, X_2, \dots, X_n$ be iid Gaussian with unity variance and unknown mean θ .
 Further θ is known to be a 0-mean Gaussian with variance σ_θ^2 . Find the MAP estimator for θ .

Solution: We are given

$$f(\theta) = \frac{1}{\sqrt{2\pi\sigma_\theta^2}} e^{-\frac{\theta^2}{2\sigma_\theta^2}}$$

$$f(x|\theta) = \frac{1}{(\sqrt{2\pi})^n} e^{-\frac{\sum_{i=1}^n (x_i - \theta)^2}{2}}$$

❖ MAP equation is

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

We will consider one example let $X_1, X_2, \dots, x_n$ be iid Gaussian with unit variance and unknown mean θ further θ is known to be a zero mean Gaussian with a variance σ_θ^2 this is the prior PDF that is Gaussian with variance σ_θ^2 , find the map estimator for

theta? We are given f theta this is the PDF prior PDF $\frac{1}{\sqrt{2\pi\sigma_\theta^2}} e^{-\frac{\theta^2}{2\sigma_\theta^2}}$ into e to the power minus theta square divided by 2 Sigma Theta square.

So this is the 0 mean Gaussian with variance sigma theta square and the likelihood function f of x given theta that is also Gaussian with unknown mean theta, so that way f of X given theta is equal to one over two pi to the power n into e to the power minus summation, I going from 1 to N of X I minus theta whole square divided by 2. This is the likelihood function of random data. Now we will consider the map equation del Delta of log F theta plus del del theta log fx given theta at theta hat equal to theta hat map must be equal to 0.

(Refer Slide Time: 16:52)

Example....

$$\diamond \text{We get } \frac{\partial}{\partial \theta} \left[\log f(\theta) + \log f(x|\theta) \right]_{\theta=\hat{\theta}_{MAP}} = 0$$

$$\Rightarrow \hat{\theta}_{MAP} = \frac{1}{N+1} \sum_{i=1}^N x_i$$

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

Two observations: (1) As $\sigma_\theta^2 \rightarrow \infty$, $\hat{\theta}_{MAP} \rightarrow \hat{\theta}_{MLE}$

(2) As $N \rightarrow \infty$, $\hat{\theta}_{MAP} \rightarrow \hat{\theta}_{MLE}$

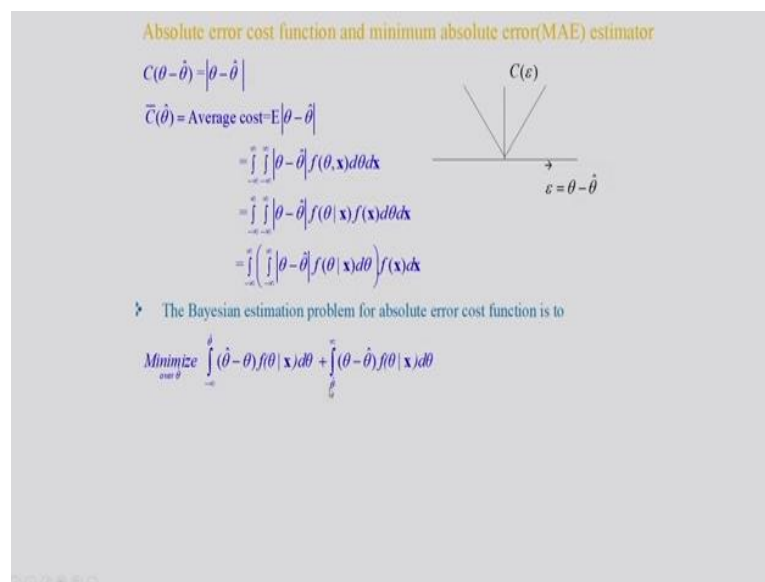
And this will give us this relationship theta by sigma theta square minus summation X I minus theta I going from 1 to N F when theta equal to theta hat map is equal to 0. So this will give theta hat map is equal to 1 by n plus 1 by sigma theta square into summation X i i going from 1 to n this is the expression for theta hat map. And now we can manipulate this expression if we divide by n then multiplied by n therefore n divided by n plus 1 by sigma theta squared into 1 by n into summation X i i going from 1 to n.

Now this quantity is theta hat MLE, so therefore theta hat map will be equal to n divided by n plus 1 by sigma theta square into theta hat MLE, it is a scaled version of theta hat MLE. Two observations are made number 1 as sigma theta square goes to infinity then this quantity will

become 0 and n by n will be equal to 1, so in the case $\hat{\theta}$ will approach to θ hat MLE.

If σ^2 that is the variance of the prior PDF is very large in that case $\hat{\theta}$ will be approximately θ hat MLE. Now next is if number of data point is large, so as n tends to infinity, so in that case this ratio will approach 1 therefore $\hat{\theta}$ will be again I proceed to θ hat MLE. So as n becomes lost did I add map will tend to θ hat MLE and as σ^2 become large then also $\hat{\theta}$ will tend to θ hat MLE.

(Refer Slide Time: 19:01)



And now we will briefly introduce absolute error cost function and minimum absolute error MAE estimator. C of θ minus $\hat{\theta}$ is equal to absolute value of θ minus $\hat{\theta}$ this is the cost function, so as a function of error this is given by and positive a pair already positive this is a positive slope and if at all is negative then this is slope is negative. so average cost now $\bar{C}$ up $\hat{\theta}$, that is the Bayesian risk or ever average cost function.

That is equal to E of absolute value of θ minus $\hat{\theta}$ and that is double integration from minus infinity to infinity minus infinity to infinity mod of θ minus $\hat{\theta}$ into F of θ , X D θ D X . This is the joint PDF between θ and X . Now writing in terms of prior and conditional PDF we write this double integral as double integral of mod of θ minus $\hat{\theta}$ into F of θ given X into F X D θ dx .

As earlier we can now take $\int_{-\infty}^{\infty} f(\theta|x) d\theta$ outside so that we have now two integrals one outer integral and one inner integral and this quantity is always greater than equal to 0, therefore the minimum can be obtained by minimizing this expression only. Thus the Bayesian estimation problem for absolute error cost function is to minimize the word $\int_{-\infty}^{\infty} f(\theta|x) d\theta$ from minus infinity to $\theta - \hat{\theta}$ this is one part.

And the other part is $\int_{\theta - \hat{\theta}}^{\infty} f(\theta|x) d\theta$ into $\int_{-\infty}^{\infty} f(\theta|x) d\theta$ in this interval $\int_{-\infty}^{\theta - \hat{\theta}} f(\theta|x) d\theta$ will be a positive quantity into $\int_{-\infty}^{\infty} f(\theta|x) d\theta$ given $X = D$, so this is the integration and other side $\int_{\theta - \hat{\theta}}^{\infty} f(\theta|x) d\theta$ will be negative therefore we write here as $\int_{-\infty}^{\theta - \hat{\theta}} f(\theta|x) d\theta$ given $X = D$. So that way we have integration from minus infinity to $\theta - \hat{\theta}$ and $\theta - \hat{\theta}$ to infinity.

(Refer Slide Time: 21:45)

Minimum absolute error (MAE) estimator

❖ Solving the estimation, we can show that $\hat{\theta}_{MAE}$ satisfies

$$\int_{-\infty}^{\hat{\theta}_{MAE}} f(\theta|x) d\theta = \int_{\hat{\theta}_{MAE}}^{\infty} f(\theta|x) d\theta = \frac{1}{2}$$

❖ Thus, $\hat{\theta}_{MAE}$ is the median of the *a posteriori* PDF.

Solving the estimation problem we can show that $\theta - \hat{\theta}_{MAE}$ satisfy this equation that is integration from minus infinity to $\theta - \hat{\theta}_{MAE}$ of $f(\theta|x)$ given $X = D$ is equal to integration from $\theta - \hat{\theta}_{MAE}$ to infinity of $f(\theta|x)$ given $X = D$ and this is equal to $\frac{1}{2}$. Thus $\theta - \hat{\theta}_{MAE}$ correspond to the median of the a-posteriori PDF. So we have to find out the median of the a posteriori PDF this is the approach to a PDF its median value is $\theta - \hat{\theta}_{MAE}$.

(Refer Slide Time: 22:32)

Note that

- ❖ $\hat{\theta}_{MMSE}$ is the mean of $f(\theta | \mathbf{x})$
- ❖ $\hat{\theta}_{MAP}$ is the mode of $f(\theta | \mathbf{x})$
- ❖ $\hat{\theta}_{MAE}$ is the median of $f(\theta | \mathbf{x})$

For symmetric PDFs, these estimators coincide.

Now we know that $\hat{\theta}_{MMSE}$ is the conditional mean it is mean of f of θ given $\mathbf{X}$ to that $\hat{\theta}_{MAP}$ is the mode of f of θ given $\mathbf{X}$, and finally $\hat{\theta}_{MAE}$ is the median of f of θ given $\mathbf{X}$. So $\hat{\theta}_{MMSE}$ is given by the mean of the a-posteriori PDF $\hat{\theta}_{MAP}$ is given by the mode of a posteriori PDF and that $\hat{\theta}_{MAE}$ is the median of a-posteriori PDF. for symmetric PDFs these estimators will coincide portion because symmetric PDFs mean more than median coincide.

(Refer Slide Time: 23:21)

Summary

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

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

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

Applying Bayes rule, the MAP equation can be written as

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

Let us summarize the map estimator minimizes the hit or miss cost function and is given by $\hat{\theta}_{MAP}$ is equal to $\arg \max_{\theta} f(\theta | \mathbf{x})$. So a-posteriori PDF we have to consider and the mode of the a-posteriori distribution will give to that $\hat{\theta}_{MAP}$ is the mode of

the a posterity PDF. If f of θ given X is differentiable then θ hat map is given by the map equations that is either in terms of the PDF $\text{del Del } \theta$ of $F \theta$ given X map is equal to 0.

Or in terms of the log-likelihood function, $\text{del Del } \theta$ of $L \theta$ given X θ hat map is equal to 0. Applying Bayes rule the map equations can be written in this form $\text{del Del } \theta$ of $\log$ of $f \theta$ at θ ab plus $\text{del Del } \theta$ of $\log$ of $f x$ given θ at θ ab must be equal to 0.

(Refer Slide Time: 24:39)

Summary...

- ❖ The MAP estimator is related to the MLE. When $f(\theta)$ is sufficiently flat, $\ln(f(\theta))$ will also be flat so that $\frac{\partial \ln(f(\theta))}{\partial \theta} \approx 0$.
In that case $\hat{\theta}_{MAP} \approx \hat{\theta}_{MLE}$.
- ❖ The MAE estimator minimizes the mean absolute error cost function and is given by $\hat{\theta}_{MAE} = \text{median of posterior } f(\theta | \mathbf{x})$.
- ❖ For symmetric PDFs, $\hat{\theta}_{MMSE}$, $\hat{\theta}_{MAP}$ and $\hat{\theta}_{MAE}$ coincide.

The map estimator is related to the MLE, when $F \theta$ is sufficiently flat $\log$ of $F \theta$ will also will flat, so that the partial derivative of the log likelihood function with respect to θ will be approximately equal to 0. So in that case θ hat map will be equal to θ hat MLE the minimum absolute error estimator MAE estimator minimizes this the mean absolute error cost function and is given by θ hat MAE is equal to median of the posterior PDF f of θ given X . For symmetric PDFs θ hat MMSC θ hat map and θ hat MAE coincide thank you.