

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

Lecture – 17
Optimal Linear Filters: Wiener Filter

(Refer Slide Time: 00:43)

Let us recall

- ❖ The ML principle of parameter estimation assumes the parameter θ to be an unknown constant and the estimator is given by $\hat{\theta}_{MLE} = \arg \max_{\theta} f(\mathbf{x}; \theta)$
- ❖ Bayesian estimators assume the parameter θ to be a random variable with the prior PDF $f(\theta)$ or a prior PMF $p(\theta)$.
- ❖ A Bayesian estimator associates a cost function to the estimation error and designs the estimator by minimizing the average of the cost function.
- ❖ Depending on the cost function considered, we established
 - $\hat{\theta}_{MMSE}$ 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 $f(\theta | \mathbf{x})$

Hello students, welcome to the lecture on optimal linear filters; let us recall the ML principle maximum likelihood principle of parameter estimation assumes the parameter theta to be an unknown constant and the estimator is given by theta at MLE equal to arg max theta of fx, theta, this is the likelihood function. Bayesian estimators assume the parameter theta to be a random variable with the prior PDF f theta or a prior PMF p theta.

A Bayesian estimator associates a cost function to the estimation error and designs the estimator by minimizing the average of the cost function, depending on the cost function considered we established theta hat MMSE that is the mean of the posterior PDF f theta given x, theta hat MAP, maximum a posteriori probability, theta hat MAP is the mode of the posterior PDF f theta given x, theta hat MAE is the median of the posterior PDF f theta given x, MAE here stands for minimum absolute error.

(Refer Slide Time: 02:11)

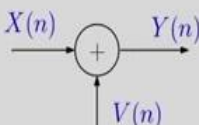
- ❖ Now we will explore extending these parameter estimation principles to estimate signals from noisy data.
- ❖ We will introduce the powerful Wiener filtering technique in this lecture.

Now, we will explore extending these parameter estimation principles to estimate signals from noisy data, we will introduce the powerful Wiener filtering technique in this lecture.

(Refer Slide Time: 02:26)

Estimation of signal in presence of white Gaussian noise (WGN)

- ❖ Consider the signal model is $Y(n) = X(n) + V(n)$



where $Y(n)$ is the observed signal, $X(n)$ is the original signal and $V[n]$ is a white Gaussian with mean 0 and variance σ_v^2 .

- ❖ The problem is to recover $X(n)$ from the noisy observations $Y(n)$.
- ❖ This is the classical signal denoising problem.
- ❖ The classical frequency selective filters cannot be applied to filter white noise, because it is spread over the entire frequency band.

Let us start with the problem of estimation of signal in the presence of white Gaussian, this is abbreviated as WGN. Consider the signal model $Y_n = X_n + V_n$, this is the addition X_n plus V_n , we are getting Y_n , where Y_n is the observed signal, X_n is the original signal and V_n is a white Gaussian noise with mean 0 and variance σ_v^2 . Now, our problem is to recover X_n from the noisy observation Y_n , this is noisy observation.

From noisy observation, we want to recover X_n , so this is the classical signal denoising problem because this signal is noisy; we want to remove the noise. The classical frequency selective filters cannot be applied to filter white noise because it is spread over the entire

frequency band because we know that white noise has uniform power spectral density over the entire frequency band, indicate some discrete time white noise, the power spectrum is spread from minus pi to pi.

(Refer Slide Time: 04:00)

ML Estimation of $X(n)$

- ❖ Maximum likelihood estimation for $X(n)$ determines that value of $X(n)$ for which the random samples $Y(i), i = 1, 2, \dots, n$ are the most likely.
- ❖ Let us represent the random samples $Y(i), i = 1, 2, \dots, n$ by $\mathbf{Y}(n) = [Y(n) \ Y(n-1) \dots Y(1)]^T$ and the particular values, $y(1), y(2), \dots, y(n)$ by $\mathbf{y}(n) = [y(n) \ y(n-1) \dots y(1)]^T$. Further assume the $Y(i)$ s to be i.i.d.
- ❖ The likelihood function $f(\mathbf{y}(n) / x(n))$ will be Gaussian with mean $x(n)$

$$f(\mathbf{y}(n) / x(n)) = \frac{1}{(\sqrt{2\pi})^n} e^{-\sum_{i=1}^n \frac{(y(i) - x(n))^2}{2\sigma_y^2}}$$

Therefore, the classical frequency selective filters cannot be applied to filter white noise, let us start with ML estimation of X_n , maximum likelihood estimation for X_n determines that value of X_n for which random samples Y_i ; i going from 1 to n are the most likely. Let us represent a random samples Y_i ; i going from 1 to n , we will represent it by this observed data vector, $\mathbf{Y}_n$ vector is equal to Y_n, Y_{n-1} up to Y_1 transpose.

This is the way we represent data $\mathbf{Y}_n$ vector and the particular values of this random vector are given by y_1, y_2 up to y_n and it is denoted by $\mathbf{y}_n$ vector that is the vector comprising of y_n, y_{n-1} up to y_1 transpose. Further, we assume that Y_i is to be iid, to apply the ML estimation criterion, we assumed the Y_i 's to be iid; independent and identically distributed. Now, the likelihood function $f(\mathbf{y}_n)$ given x_n will be Gaussian with mean x_n .

So, we can write $f(\mathbf{y}_n)$ given x_n is equal to 1 over $(\sqrt{2\pi})^n$ into e to the power minus summation $(y_i - x_n)^2$ divided by $2\sigma_y^2$, i going from 1 to n , so this we get assuming x_i is a iid, therefore each is Gaussian distributed, therefore they will be jointly Gaussian distributed like this.

(Refer Slide Time: 06:01)

ML Estimation of $X(n)$...

❖ $\hat{x}_{MLE}(n)$ is given by

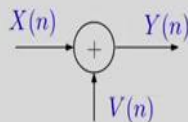
$$\frac{\partial}{\partial x(n)} (f(y(1), y(2), \dots, y(n)) / x[n])|_{\hat{x}_{MLE}(n)} = 0$$
$$\Rightarrow \hat{x}_{MLE}(n) = \frac{1}{n} \sum_{i=1}^n y(i)$$

And now except MLE is given by the partial derivative of the likelihood function with respect to X_n at $\hat{X}_{MLE}$ is equal to 0 and from this equation, we get $\hat{X}_{MLE}$ and is equal to $\frac{1}{n} \sum_{i=1}^n y_i$, i going from 1 to n , so this way we can estimate the signal value from a set of observed data.

(Refer Slide Time: 06:33)

Bayesian estimation of $X(n)$

❖ Assume $X(n)$ to be random. We can apply the MMSE, MAP and MAE estimation.



Suppose $X(n)$ is 0-mean Gaussian with variance σ_x^2 and $V[n]$ is a white Gaussian with mean 0 and variance σ_v^2 .

❖ Assume both σ_x^2 and σ_v^2 to be known. The problem is to find the best guess for $X(n)$ given the observation $Y(i), i=1, 2, \dots, n$.

Next, we consider the Bayesian estimation of X_n , here we will assume X_n to be random, we can apply MMSE minimum mean square error estimation MAP; maximum a posteriori probability, MAE minimum absolute error estimation techniques, so this is the model as earlier; Y_n is $X_n + V_n$, this V_n is white Gaussian noise and X_n has some prior probability distribution, suppose X_n is 0 mean Gaussian with variance σ_x^2 .

And V_n is a white Gaussian noise with mean 0 and variance σ_v^2 , so we assume X_n to be 0 mean Gaussian with variance σ_x^2 and V_n to be Gaussian white noise with mean 0 and variance σ_v^2 . Assume both σ_x^2 and σ_v^2 to be known, so we assume that both are known samples. The problem is to find the best guess for X_n given the observation Y_i ; i going from 1 to n .

(Refer Slide Time: 07:56)

Bayesian estimation of $X(n)$...

❖ To find $\hat{x}_{MAP}(n)$ and $\hat{x}_{MMSE}(n)$, we have to find *a posteriori* PDF

$$f(x(n)/y(n)) = \frac{f(x(n))f(y(n)/x(n))}{f(y(n))}$$

$$= \frac{1}{f(y(n))} e^{-\frac{1}{2\sigma_x^2}x^2(n) - \sum_{i=1}^n \frac{(y(i)-x(n))^2}{2\sigma_v^2}}$$

❖ Taking the logarithm on both sides, we get

$$\log_e f(x(n)/y(n)) = -\frac{1}{2\sigma_x^2}x^2(n) - \sum_{i=1}^n \frac{(y(i)-x(n))^2}{2\sigma_v^2} - \log_e f(y(n))$$

So, this is the Bayesian estimation problem, to find $\hat{x}_{MAP}$ and $\hat{x}_{MMSE}$, we have to find out the a posteriori PDF that is f of x_n , given y_n , y_n is a vector that is f of x_n into f of y_n given x_n divided by f of y_n , so this we can write as; because this part is e to the power minus 1 by 2 σ_x^2 into $x^2(n)$ and this part is given by e to the power minus summation $y_i - x_n$ whole square divided by 2 σ_v^2 , i going from 1 to n .

So, we can write the posterior PDF by this relationship where $f(y_n)$ does not involve x_n , now we will take the logarithm on both sides, so $\log$ of the a posteriori PDF that is equal to minus 1 by 2 σ_x^2 into $x^2(n)$ minus summation, i going from 1 to n of $y_i - x_n$ whole square divided by 2 σ_v^2 minus $\log$ of $f(y_n)$. So, now we can apply the condition for MAP and MMSE.

(Refer Slide Time: 09:19)

Bayesian estimation of $X(n)$...

- ❖ $\log_e f(x(n)/\mathbf{y}(n))$ is maximum at $\hat{x}_{MAP}(n)$. Therefore, taking partial derivative of $\log_e f(x(n)/\mathbf{y}(n))$ with respect to $x(n)$ and equating it to 0, we get

$$\left. \frac{x(n)}{\sigma_x^2} - \sum_{i=1}^n \frac{(y(i) - x(n))}{\sigma_v^2} \right|_{\hat{x}_{MAP}(n)} = 0$$

$$\hat{x}_{MAP}(n) = \frac{\sum_{i=1}^n y(i)}{n + \frac{\sigma_v^2}{\sigma_x^2}}$$

- ❖ Similarly, the MMSE estimator is given by

$$\hat{x}_{MMSE}(n) = E(X(n)/\mathbf{y}(n)) = \frac{\sum_{i=1}^n y(i)}{n + \frac{\sigma_v^2}{\sigma_x^2}}$$

Now, this log of the a posterior PDF is maximum at $\hat{x}_{MAP} n$, therefore taking partial derivative of log of f_x given n with respect to x_n and equating it to 0, we get this condition; x_n divided by σ_x^2 minus summation, i going from 1 to n of $y_n - x_n$ divided by σ_v^2 $\hat{x}_{MAP} n$ is equal to 0, so this is the condition we get by applying partial derivative with respect to x_n .

And if we simplify this expression, we will get $\hat{x}_{MAP} n$ that is the maximum a posterior estimate of x_n as summation of y_i ; i going from 1 to n divided by $n + \sigma_v^2$ divided by σ_x^2 similarly, the MMSE is given by $\hat{x}_{MMSE} n$ is equal to conditional mean of the posterior PDF that is E of X_n given y_n , so that is given by summation y_i ; i going from 1 to n divided by $n + \sigma_v^2$ by σ_x^2 .

(Refer Slide Time: 10:50)

WSS assumption

- ❖ In the previous example, we assumed $Y(i)$ s to be iid. The signal samples are generally dependent and the general case of signal estimation is difficult. We look for a simpler estimator.
- ❖ A simpler class of MMSE estimators, known as the *linear MMSE (LMMSE) estimators*, are available for WSS signals. It exploits the joint correlation structure of $X(n)$ and $Y(n)$.
- ❖ As $X(n)$ and $Y(n)$ are zero-mean, they are characterized by
 - the autocorrelation function $R_X(m) = EX(n)X(n+m)$
 - the cross-correlation between $X(n)$ and $Y(n)$ vector

$$\mathbf{r}_{XY} = EX(n)\mathbf{Y}(n) = \begin{bmatrix} EX(n)Y(n) \\ EX(n)Y(n-1) \\ \vdots \\ EX(n)Y(1) \end{bmatrix} = \begin{bmatrix} R_{XY}(0) \\ R_{XY}(1) \\ \vdots \\ R_{XY}(n) \end{bmatrix}$$

Actually, the same quantity we will get because of the symmetry of the Gaussian PDF, in the previous example we assumed Y_i is to be iid, the signal samples are generally dependent, so iid assumption we cannot use and the general case of signal estimation is difficult, we look for a simple estimator. A simple class of MMSE estimators known as the linear MMSE; LMMSE estimators are available for WSS signals.

We have to consider the signals to be white sense stationary, it exploits the joint correlation structure of X_n and Y_n in terms of autocorrelation function and cross correlation. We are assuming X_n, Y_n to be 0 mean therefore, they are characterized by the autocorrelation function R_x of m that is E of X_n into X of $n + m$, the cross correlation between X_n and Y_n vector is given by small r_{xy} vector, this is the notation and this is equal to E of X_n into Y_n vector.

And this is given by this expression E of X_n into Y_n , E of X_n into Y_{n-1} and so on up to E of X_n into Y_1 and using the joint stationarity property, we get this term is equal to r_{xy} of 0, this term is equal to r_{xy} of 1 and this term is equal to r_{xy} of $n-1$.

(Refer Slide Time: 12:31)

▪ the autocorrelation matrix of random vector $\mathbf{Y}(n)$

$$\mathbf{R}_Y = E\mathbf{Y}(n)\mathbf{Y}'(n) = \begin{bmatrix} R_Y(0) & R_Y(1) & R_Y(n-1) \\ R_Y(1) & R_Y(0) & R_Y(n-2) \\ \vdots & \vdots & \vdots \\ R_Y(n-1) & R_Y(n-2) & R_Y(0) \end{bmatrix}$$

The autocorrelation matrix of the random vector Y_n is given by this R_Y matrix that is equal to E of Y_n into Y_n transpose, this will be a matrix and we will take the expected of the individual elements and assuming stationarity we get R_{Y0}, R_{Y1} up to R_{Yn-1} similarly, second row is R_{Y1}, R_{Y0} and this way up to R_{Yn-2} and so on, last row will be R_{Yn-1}, R_{Yn-2} up to R_{Y0} , so this is the autocorrelation matrix.

(Refer Slide Time: 13:14)

Jointly Gaussian case

❖ Suppose $X(n)$ and $\mathbf{Y}(n)$ are zero-mean and jointly Gaussian. Then,

$$X(n) \sim N(0, R_X(0)) \text{ and } \mathbf{Y}(n) \sim N(0, \mathbf{R}_Y)$$

$$\hat{X}_{MMSE}(n) = E[X(n) | \mathbf{Y}(n)]$$

It can be shown that for the jointly Gaussian case

$$\begin{aligned} \hat{X}_{MMSE}(n) &= E[X(n) | \mathbf{Y}(n)] \\ &= (\mathbf{R}_Y^{-1} \mathbf{r}_{XY})' \mathbf{Y}(n) = \mathbf{r}_{XY}' \mathbf{R}_Y^{-1} \mathbf{Y}(n) \end{aligned}$$

which is a linear function of $Y(n), Y(n-1), \dots, Y(1)$

❖ In other words, $X(n)$ and $\mathbf{Y}(n)$ are zero-mean and jointly Gaussian, $X(n)$

can be estimated using a linear filter $\mathbf{h}(n) = [h(0) \ h(1) \ \dots \ h(n)]'$

$$\mathbf{h}(n) = \mathbf{R}_Y^{-1} \mathbf{r}_{XY}$$

We will consider a simple case that is the jointly Gaussian case suppose, X_n and Y_n are 0 mean and jointly Gaussian, then X_n is distributed as normal with mean 0 and variance R_X of 0 because 0 mean therefore, variance will be R_X of 0 only and Y_n is distributed as normal with mean 0 vector and covariance R_Y matrix, so again this covariance matrix is equal to R_Y because of zero mean.

In that case, $\hat{X}_{MMSE}$ at point n will be given by E of X_n given Y_n vector, it can be shown that for jointly Gaussian case, this MMSE; $\hat{X}_{MMSE}$ is given by E of X_n , given Y_n that is equal to transpose of R_Y inverse into r_{xy} into Y_n vector. Now, writing this transpose as r_{xy} vector into R_Y inverse, we are taking the transpose of the product, so we will get like this.

And R_Y is anyway is a symmetric matrix therefore this will remain same, so we will get r_{xy} vector transpose into R_Y inverse into Y_n clearly, this is a linear combination of Y_n, Y_{n-1} up to Y_1 , in other words if X_n and Y_n are 0 mean and jointly Gaussian, then X_n can be estimated using a linear filter h_n which is given by a vector comprising of h_0, h_1 up to h_{n-1} .

And it is related to the autocorrelation matrix and cross covariance vector by this relationship; h_n is equal to R_Y inverse into r_{xy} vector.

(Refer Slide Time: 15:32)

Optimal filtering or Wiener filtering

- ❖ Inspired by the Gaussian case, a mathematically simple and computationally easier estimator is obtained by assuming a linear filter structure for the estimator.
- ❖ The minimization of errors then results in the determination of an optimal set of filter parameters using the MMSE.
- ❖ Linearity assumption makes the estimation problem technically simple, because linear filters are easily implemented. The problem becomes mathematically simple, because of the resulting quadratic optimization problem.

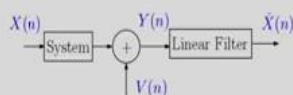
Inspired by the Gaussian case, a mathematically simple and computationally easier estimator is obtained by assuming a linear filter structure for the estimator. Now, for the estimator we will assume a linear filter structure, the minimization of errors then results the determination of an optimal set of filter parameters using the MMSE principle. Now, why linear assumption; linearity assumption makes the estimation problem technically simple, why?

Because linear filters are easily implemented, it is simply multiplication and summation, the problem becomes mathematically simple because the resulting quadratic optimization problem because of this linear filter structure, the mean square error criterion will be an a quadratic optimization problem.

(Refer Slide Time: 16:36)

Linear Minimum Mean Square Error (LMMSE) Estimator

- ❖ The filtering principle can be explained with the help of the following block diagram



- ❖ The filter operates on a block of random samples where M and N are constant and may be infinite

$$\begin{array}{ccccccc} Y(n-M+1) & \dots & Y(n) & \dots & Y(n+N) \\ \hline & & n-M+1 & & n & & n+N \end{array}$$

- ❖ The linear filter is characterised by a set of filter parameters $h(-N), h(-N+1), \dots, h(0), h(1), \dots, h(M-1)$ such that

$$\hat{X}(n) = \sum_{i=-N}^{M-1} h(i)Y(n-i)$$

Now, we will discuss the linear minimum mean square error LMMSE estimator principle; the filtering principle can be explained with the help of the block diagram here, so suppose X_n it is pass through some system and the output is also corrupted by noise V_n and this is my observed data Y_n and we have to apply some linear filter to get back $\hat{X}_n$, which is an estimate of X_n .

Now, this linear filter operates on a block of random samples where M and N are constant and may be infinite, so this is the block of data, this start from $n - M + 1$; Big M , this is n and this is $n + \text{capital } N$, we have both past and future data, so it starts at point $n - M + 1$ and end at $n + \text{capital } N$ where M and N may be infinite. The linear filter is characterized by a set of filter parameters corresponding to suppose, this block of data h of $-N$, h of $-N + 1$, then up to h_0 , then h_1 up to h of $M - 1$.

Such that now, filtered output will be given by $\hat{X}_n$ that is equal to summation h_i into Y_{n-i} ; i going from $-N$ to $M - 1$, so we can get $\hat{X}_n$ the estimator for X_n by this relationship that is the convolution operation which is given by summation h_i into y_{n-i} ; i going from $-N$ to $M - 1$.

(Refer Slide Time: 18:38)

Linear Minimum Mean Square Error (LMMSE) principle

❖ The estimation problem can be stated as follows:
 Given random observations $Y(n-M+1), Y(n-M+2), \dots, Y(n), \dots, Y(n+N)$,
 determine an optimal set of parameters $h(-N), h(-N+1), \dots, h(0), h(1), \dots, h(M-1)$
 such that

$$\hat{X}(n) = \sum_{i=-N}^{M-1} h(i)Y(n-i)$$

and the mean square error $E\{X(n) - \hat{X}(n)\}^2$ is a minimum.

❖ Thus, we have to minimize the MSE

$$E\{X(n) - \sum_{i=-N}^{M-1} h(i)Y(n-i)\}^2$$

with respect to $h(-N), h(-N+1), \dots, h(0), h(1), \dots, h(M-1)$

Now, the estimation problem can be stated as follows; given random observations Y_{n-M+1} , Y_{n-M+2} and so on up to Y_n , then up to $Y_{n+\text{capital } N}$, determine an optimal set of parameters, these parameters are h of $-N$, h of $-N + 1$ up to h_0 , h_1 , up to h of $M - 1$, such that that estimator is the filter output $\hat{X}_n$ is equal to summation $h_i Y_{n-i}$; i going from $-N$ to $M - 1$ and the mean square error you have $X_n - \hat{X}_n$ whole square is a minimum. So,

we have the estimator output at the convolution between h_n and Y_n sequence such that the mean square error E of $X_n - \hat{X}_n$ whole square is a minimum.

Thus we have to minimize the MSE mean square error, what is the mean square error; E of X_n minus this is the estimator summation $\sum_{i=-N}^{M-1} h_i Y_{n-i}$ whole square, so this is the mean square error this we have to minimize, with respect to the filter parameters h of $-N$, h of $-N+1$ up to h_0 , h_1 up to h of $M-1$, so we have to minimize this mean square error with respect to these parameters.

(Refer Slide Time: 20:23)

LMMSE estimation problem

- ❖ Thus, the LMMSE estimation problem is

$$\text{Minimize } E \left(X(n) - \sum_{i=-N}^{M-1} h(i)Y(n-i) \right)^2 \text{ over } h(i), i = -N \text{ to } M-1$$

- ❖ Note that the above is a quadratic optimization problem in terms of $h(i)$ s. Therefore, a unique minimum exists.
- ❖ This minimization problem results in an elegant solution if we assume jointly wide-sense stationarity of the signals $X(n)$ and $Y(n)$. The estimator parameters can be obtained from the second order statistics of the processes $X(n)$ and $Y(n)$.
- ❖ The problem of determining the estimator parameters by the LMMSE criterion is also called the *Wiener filtering problem*. The Wiener filter theories are due to Wiener and Kolmogorov.

So, we can write the LMMSE estimator some problem as minimize E of $X_n - \sum_{i=-N}^{M-1} h_i Y_{n-i}$ whole square over h_i ; i going from $-N$ to $M-1$. Now, let us see this optimization problem, this is the cost function and this cost function is a quadratic cost function in terms of h_i 's, therefore the above is a quadratic optimization problem in terms of h_i 's.

Therefore, a unique minimum exists because it is a quadratic optimization problem, unique optimum exists, the minimization problem results in an elegant solution if we assume jointly wide-sense stationarity of the signals X_n and Y_n , this we will assume that signal n observed data are jointly WSS, the estimator parameters can be obtained from the second order statistics that is autocorrelation functions and cross correlation functions of the process X_n and Y_n .

The problem of determining the estimator parameters by the LMMSE criterion is also called the Wiener filtering problem. The Wiener filter theories are due to Wiener and Kolmogorov.

(Refer Slide Time: 21:59)

Subclass of Wiener filtering problem

- ❖ Three subclasses of the problem are identified
 1. The optimal smoothing problem $N > 0$
 2. The optimal filtering problem $N = 0$
 3. The optimal prediction problem $N < 0$

$$\begin{array}{ccccccc}
 Y(n-M+1) & \dots & Y(n) & \dots & Y(n+N) \\
 \hline
 n-M+1 & & n & & n+N
 \end{array}$$

- ❖ In the smoothing problem, an estimate of the signal is made at a location inside the observation window.
- ❖ The filtering problem estimates the current value of the signal on the basis of the present and past observations.
- ❖ The prediction problem addresses the issues of optimal prediction of the future value of the signal on the basis of present and past observations.

There are three subclasses of the Wiener filtering problem that is optimal smoothing problem, if N is greater than 0; capital N is greater than 0, the optimal filtering problem if N is 0 and the optimal prediction problem if N is less than 0 that means, this N ; $n + N$, suppose if N is greater than 0 then we will consider the future samples also and therefore this will be this smoothing problem.

In smoothing problem, an estimate of the signal is made at the location inside observation window because observation window is up to here and we are making an estimator for this signal at this instant. The filtering problem estimates the current value of the signal on the basis of the present and past observation because if big N is equal to 0, then data will considering up to this point therefore, we will be considering the present and the past observations.

Now, for the prediction problem, N is less than 0 that means, we considered only past data because our begin is this side therefore, we will consider only the past data therefore, the prediction problem addresses the issues of optimal prediction of the future value of the signal on the basis of present and the past observations, these are the 3 subclasses of Wiener filtering problem.

(Refer Slide Time: 23:34)

Wiener-Hopf Equations

- ❖ The mean-square error of estimation is given by

$$Ee^2(n) = E(X(n) - \hat{X}(n))^2$$

$$= E\left(X(n) - \sum_{i=-N}^{M-1} h(i)Y(n-i)\right)^2$$

We have to minimize $Ee^2[n]$ with respect to each $h[i]$ to get the optimal estimation.

- ❖ At the minimum,

$$\frac{\partial Ee^2(n)}{\partial h(j)} = 0, \text{ for } j = -N, 0, \dots, M-1$$

(E and $\frac{\partial}{\partial h(j)}$ can be interchanged)

$$Ee(n)Y(n-j) = 0, \quad j = -N, \dots, 0, 1, \dots, M-1$$

Now, the mean square error of estimation is given by E of e square n that is equal to E of $X_n - \hat{X}_n$ whole square and if we write $\hat{X}_n$ as the linear combination of h_i and Y_{n-i} 's, we will get the same expression is equal to E of $X_n - h_i$ into Y_{n-i} ; i going from $-N$ to $M-1$ whole square, so this is the mean square error. We have to minimize E of e square n with respect to h_i 's to get the optimal estimation.

Now, at the minimum the partial derivative of E of e square n with respect to all h_j is equal to 0, so $\frac{\partial}{\partial h(j)} E$ of e square n is equal to 0 for j is equal to $-N$ up to $M-1$ now, this $\frac{\partial}{\partial h(j)}$ operation and E operation can be interchanged, so first we will take the partial derivative of E square n that is E of n and then partial derivative of E_n with respect to h , so that way we will get only Y of $n-j$.

Because if I take the partial derivative of this error term with respect to h_j , there is only one term, h_j into Y_{n-j} which involve h_j , therefore by taking the partial derivative we will get E of; E_n into Y_{n-j} is equal to 0 for j is equal to $-N$ up to $M-1$.

(Refer Slide Time: 25:28)

Wiener-Hopf Equations ...

or

$$E \left(\overbrace{X(n) - \sum_{i=-N}^{M-1} h(i)Y(n-i)}^{e[n]} \right) Y(n-j) = 0, \quad j = -N, \dots, 0, 1, \dots, M-1$$

$$R_{XY}(j) = \sum_{i=-N}^{M-1} h(i)R_Y(j-i), \quad j = -N, \dots, 0, 1, \dots, M-1$$

- ❖ This set of $N + M + 1$ equations are called *Wiener-Hopf equations* or *Normal equations*.
- ❖ We have a set of linear equations. We will see how to solve this set of equations particularly when M and N are infinite.

For N we can substitute this, E of X_n minus summation $h_i Y_{n-i}$; i going from $-N$ to $M-1$ into Y_{n-j} that is equal to 0, for j equal to $-N$ to $M-1$ now, using the joint stationarity property we will get E of X_n into Y_{n-j} is equal to R_{xyj} and therefore, we have R_{xy} of j is equal to summation h_i into R_y of $j-i$; i going from $-N$ to $M-1$, where j takes the values; $-N, -N+1$ up to $M-1$. Now, we get a set of $N + M + 1$ equations called Wiener Hopf equations or normal equations.

So, this set is the Wiener Hopf equation, we have a set of linear equation we will see how to solve this set of equations particularly, when M and N are infinite, hence such solution of linear equations are easy but here problem is both N and M may become infinite, how to solve the Wiener Hopf equation in those situation is a problematic case.

(Refer Slide Time: 26:56)

Summary

- ❖ The signal estimation problem involves estimating a signal $x(n)$ from noisy observations $Y(i)$ s using the model

$$Y(n) = X(n) + V(n)$$
 where $V(n)$ is a white Gaussian noise.
- ❖ The parameter estimation principles like MLE, MMSE and MAP can be extended to signal estimation.
- ❖ $X(n)$ and $Y(n)$ are zero-mean and jointly Gaussian, MMSE estimator for $X(n)$ results in a linear filter
- ❖ A mathematically simple and computationally easier estimator is obtained by assuming a linear filter structure for the estimator.

Let us summarize; the signal estimation problem involves estimating a signal x_n from noisy observations Y_i 's using the model Y_n is equal to $X_n + V_n$, where V_n is a white Gaussian noise. The parameter estimation principles like MLE, MMSE and MAP can be extended to signal estimation case, if X_n and Y_n are 0 mean and jointly Gaussian, then MMSE estimator for X_n results in a linear filter.

This is an important observation; if X_n and Y_n are zero mean, jointly Gaussian random processes, then MMSE estimator for X_n results in a linear filter, a mathematically simple and computationally easier estimator is obtained by assuming a linear filter structure for the estimator, so we assume a linear filter structure for the estimator.

(Refer Slide Time: 28:01)

Summary ...

- ❖ The LMMSE estimation problem can be stated as follows:
Given random observations
 $Y(n-M+1), Y(n-M+2), \dots, Y(n), \dots, Y(n+N),$
Minimize $E \left(X(n) - \sum_{i=-N}^{M-1} h(i)Y(n-i) \right)^2$ over $h(j), j = -N \text{ to } M-1$
- ❖ The minimum is given by the Wiener Hopf equations

$$R_{XY}(j) = \sum_{i=-N}^{M-1} h(i)R_Y(j-i), \quad j = -N, \dots, 0, 1, \dots, M-1$$

- ❖ This set of equations is to be solved to find the optimal filter parameters.

The LMMSE estimation problem can be stated as follows; given random observations Y_{n-M+1}, Y_{n-M+2} up to Y_{n+N} , we have to minimize E of $X_n - \sum_{i=-N}^{M-1} h_i Y_{n-i}$; i going from $-N$ to $M-1$ whole square over h_i 's; i going from $-N$ to $M-1$, so we have to minimize this mean square error with respect to h_i 's, so the minimum is given by the Wiener Hopf equations that is R_{XY} of j is equal to summation h_i into $R_Y(j-i)$; i going from $-N$ to $M-1$, where j goes from $-N$ to $M-1$. This set of equations is to be solved to find the optimal filter parameters which we will discuss in the next lecture, thank you.