

**Statistical Signal Processing**  
**Prof. Prabin Kumar Bora**  
**Department of Electronics and Electrical Engineering**  
**Indian Institute of Science – Guwahati**

**Lecture -30**  
**Adaptive Filters 4 (Contd.)**

**(Refer Slide Time: 00:38)**

- ❖ We presented the LMS adaptive filter and its variants.
  - ❖ The convergence of LMS algorithm is slow. The LLMS and NLMS algorithms have better convergence properties.
  - ❖ For hardware efficiency, the LMS algorithm is further simplified. We discussed the block LMS and sign-error LMS algorithms.
  
  - ❖ There is another class of adaptive filters with better convergence properties known as the *recursive least squares (RLS)* adaptive filters.
  - ❖ It uses the least-squares estimation principle
- Before discussing the RLS adaptive filters, we will explain the least-squares technique for parameter estimation.

Hello students. Welcome to this lecture on Least Squares Estimator. We presented the LMS adaptive filter and its variants in the previous lectures. The convergence of LMS algorithm is slow. The leaky LMS and the normalize LMS algorithms have better convergence properties. For hardware efficiency the LMS algorithm is further simplified. We discussed the block LMS and sign-error LMS algorithms.

There is another class of adaptive filters with better convergence properties. This class is known as the recursive least squares adaptive filters. It uses the least squares estimation principle. Before discussing the RLS adaptive filters we will explain the least squares techniques for parameter estimation. Least squares is a parameter estimation technique we will first examine that.

**(Refer Slide Time: 02:03)**

### Recall that

- ❖ In the parameter estimation methods discussed so far, the observed random data  $X_1, X_2, \dots, X_n$  are characterized by a known joint PDF  $f(x; \theta)$  which depends on some unobservable parameter  $\theta$
  - ❖ An MVUE has the lowest variance in the class of unbiased estimators and is the most desirable estimator.
  - ❖ We discussed the estimation approaches:
    - Method of Moments
    - Maximum likelihood approach
    - Bayesian approach.
- The LS approach is distinct from those approaches as it does not assume any probability model for the random observations.
- Oldest approach due to Gauss (1822).

Recall that in the parameter estimation methods discussed so far the observed random data  $X_1, X_2$  upto  $X_N$  are characterized by known joint PDF  $f(x; \theta)$  or a known joint PDF  $f(x; \theta)$  (02:20) which depends on some unobservable parameter  $\theta$ . The parameter  $\theta$  here is unknown and unobservable. An MVUE (Minimum Variance Unbiased Estimator) has the lowest variance in the class of unbiased estimators and is the most desirable estimator.

We discussed about various approaches to find MVUE. We discussed the estimation approaches mentioned here like method of moments, maximum likelihood approach, Bayesian approach like minimum mean square error estimation, maximum a posteriori probability estimation etcetera. The least squares approach is distinct from those approaches and it does not assume any probability model for the random observations. This is the oldest approach and due to Gauss in 1822.

**(Refer Slide Time: 03:30)**

### Least-squares Estimator

- ❖ In least square (LS) estimation, the observed data are to be a known function of some unknown parameters plus some errors which are assumed to be random.
- ❖ Unlike MVUE, MLE or the Bayesian estimators, we do not have an explicit probabilistic model for the data. The randomness is only due to the measurement error. The error is characterized by its first and second-order statistics.
- ❖ In linear least-squares estimation, a linear model is assumed for the observed data. The estimation involves minimizing sum of the squares of the observation errors.
- ❖ The linear model greatly simplifies the estimation problem.

Let us discuss about least squares estimator. In LS estimation the observed data are to be a known function of some unknown parameters plus some errors which are assumed to be random plus some noise. Unlike MVUE MLE or Bayesian estimators we do not have an explicit probabilistic model this is the distinction. The randomness is only due to the measurement error. The error is characterized by its first.

And second order statistics that is important because only mean or second order statistics like variance, covariance etcetera are used. In linear least square estimation a linear model is assumed for the observed data. We will discuss about linear least squares estimation only. The estimation involves minimizing the sum of the squares of observation errors. So estimation involves therefore here estimation means minimizing the sum of the squares of the observation errors. The linear model greatly simplifies the estimation problem. Therefore, we have a linear model and the parameters are to be estimated.

**(Refer Slide Time: 05:14)**

### Linear model and least-squares principle

- ❖ Consider a linear model where the observations  $X_1, X_2, \dots, X_N$  are represented as

$$X_i = \sum_{j=1}^K a_{ij} \theta_j + e_i, \quad i=1, 2, \dots, N$$

where  $\theta_1, \theta_2, \dots, \theta_K$  are  $K$  unknown parameters,  $a_{ij}$ s are known constants and  $e_i$ s are random measurement errors modelled as zero-mean uncorrelated RVs.

- ❖ For example, the measured positions of a particle moving with a constant acceleration can be expressed as a linear model.  $X = ut + \frac{1}{2}at^2$

- ❖ The LS estimation problem considers the **sum square error (SSE)** as the cost function given by

$$J(\theta_1, \theta_2, \dots, \theta_K) = \sum_{i=1}^N e_i^2 = \sum_{i=1}^N (X_i - \sum_{j=1}^K a_{ij} \theta_j)^2$$

$$\therefore J(\theta_1, \theta_2, \dots, \theta_K) = \sum_{i=1}^N (X_i - \sum_{j=1}^K a_{ij} \theta_j)^2$$

Let us consider a linear model where the observation  $X_1, X_2$  upto  $X_N$  are represented as  $X_i$  that = summation  $a_{ij} \theta_j$   $j$  going from 1 to  $K$  +  $e_i$  where  $i = 1, 2$  up to  $N$  this is the model of the signal. This is a linear model because the observed data is a linear function of the parameters plus some noise. Here  $\theta_1, \theta_2$  upto  $\theta_K$  are unknown parameters  $a_{ij}$ s are known constants and  $e_i$ s are random measurement errors modeled as zero-mean uncorrelated random variable.

Therefore, the error is an zero-mean random variable and its covariance of any two errors is equal to zero. Here  $\theta_1, \theta_2$  upto  $\theta_K$  are  $K$  unknown parameters  $a_{ij}$ s are known constants and  $e_i$ s are random measurement errors modeled as zero-mean uncorrelated random variables. So  $e_i$ s are zero-mean and their covariance, covariance between any two samples is equal to zero.

For example, the measured position of a particle moving with a constant acceleration can be expressed as a linear model. Suppose position  $X = ut + \frac{1}{2}at^2$  suppose this is the model then this  $u$  ( $\theta_1$ ) ( $\theta_2$ ) parameters therefore  $x$  is a linear combination of ( $\theta_1$ ) ( $\theta_2$ ). The LS estimation problem considers the sum square error. So this is SSE everywhere it as SSE so sum square error square of the errors.

And then sum up so that way this sum square error this is the cost function  $J(\theta_1, \theta_2$  upto  $\theta_K$  that is = summation  $e_i^2$  for all  $N$  observation. So that way summation  $e_i^2$   $i$  going from 1 to  $N$  that =  $X_i - \sum_{j=1}^K a_{ij} \theta_j$   $j$  going from 1 to  $K$  whole square  $i$  going from 1 to  $N$ . So this is the because  $X_i$  is given like this therefore  $e_i$  we can find out by

subtracting this term from  $X_i$ .

So that way this cost function this sum square error and it is given by this so this is the cost function  $J(\theta_1, \theta_2, \dots, \theta_K) = \sum_{i=1}^N (X_i - \sum_{j=1}^K a_{ij} \theta_j)^2$  going from 1 to K whole square i going from 1 to N. This is the cost function corresponding to LS estimation.

**(Refer Slide Time: 09:04)**

**Linear model and least-squares principle ...**

- ❖ The optimal set of parameters  $\hat{\theta}_{LS}, \hat{\theta}_{2LS}, \dots, \hat{\theta}_{KLS}$  are obtained by minimizing  $J(\theta_1, \theta_2, \dots, \theta_K)$  with respect to  $\theta_1, \theta_2, \dots, \theta_K$ .
- ❖ Thus  $\hat{\theta}_{LS}, \hat{\theta}_{2LS}, \dots, \hat{\theta}_{KLS}$  are given by
$$\hat{\theta}_{LS}, \hat{\theta}_{2LS}, \dots, \hat{\theta}_{KLS} = \arg \min_{\theta_1, \theta_2, \dots, \theta_K} J(\theta_1, \theta_2, \dots, \theta_K)$$

The optimal set of parameters that is  $\theta_1$  hat LS,  $\theta_2$  hat LS upto  $\theta_K$  hat LS are obtained by minimizing this cost function sum square error with respect to the parameters  $\theta_1, \theta_2$  upto  $\theta_K$ . Thus, we can write that  $\theta_1$  hat LS,  $\theta_2$  hat LS upto  $\theta_K$  hat LS =  $\arg \min_{\theta_1, \theta_2, \dots, \theta_K} J(\theta_1, \theta_2, \dots, \theta_K)$ . So the parameters which minimizes this cost function (()) (09:51) LS estimators.

Since our cost function is a quadric cost function we can find out the LS estimators uniquely by applying the differentiation principle. So we will take the partial derivative to J with respect to all the parameters and set them to 0.

**(Refer Slide Time: 10:20)**

### Example 1

- ❖ Suppose the observations  $X_i, i = 1, 2, \dots, N$  are related by  
$$X_i = \theta + e_i,$$
where  $\theta$  is the unknown parameter and  $e_i$  is the observation or measurement noise.

- ❖ According to LS principle, we have to minimize the sum-square error.

$$J(\theta) = \sum_{i=1}^N e_i^2 = \sum_{i=1}^N (X_i - \theta)^2 \text{ with respect to } \theta.$$

- ❖ Thus  $\hat{\theta}_{LS} = \arg \min_{\theta} J(\theta)$

$\hat{\theta}_{LS}$  is given by

$$\left. \frac{dJ(\theta)}{d\theta} \right|_{\theta=\hat{\theta}_{LS}} = 0$$

$$\Rightarrow \hat{\theta}_{LS} = \frac{1}{N} \sum_{i=1}^N X_i$$

which is same as the MLE.

We will consider one example suppose the observations  $X_i$  are related by this simple model  $X_i = \theta + e_i$  where  $\theta$  is the unknown parameter  $e_i$  is the observation or measurement noise. Now this is the sum square error  $J(\theta) = \sum_{i=1}^N (X_i - \theta)^2$  where  $i$  going from 1 to  $N$  we have to minimize this sum square error with respect to  $\theta$ . Thus,  $\hat{\theta}_{LS} = \arg \min_{\theta} J(\theta)$ .

So we have to minimize this  $J(\theta)$  with respect to  $\theta$  and it is given by  $\frac{dJ(\theta)}{d\theta} = 0$  and taking the derivative of  $J(\theta)$  with respect to  $\theta$  and setting it to 0 we get  $\hat{\theta}_{LS} = \frac{1}{N} \sum_{i=1}^N X_i$  where  $i$  going from 1 to  $N$ . So this is  $\hat{\theta}_{LS}$  which we obtain by differentiating this function and setting it to 0 and we observed that  $\hat{\theta}_{LS}$  is same as the MLE because they are in that case also we found that  $\hat{\theta}_{MLE} = \frac{1}{N} \sum_{i=1}^N X_i$  where  $i$  going from 1 to  $N$ .

**(Refer Slide Time: 12:06)**

### Example 2

❖ Suppose the observations  $X_i, i = 1, 2, \dots, N$  are related by

$$X_i = \theta_1 + a_i \theta_2 + e_i,$$

where  $\theta_1$  and  $\theta_2$  are the unknown parameters and  $e_i$  is the observation error.

❖ Then  $J(\theta_1, \theta_2) = \sum_{i=1}^N (X_i - \theta_1 - a_i \theta_2)^2$

❖ The LS estimators  $\hat{\theta}_{1LS}$  and  $\hat{\theta}_{2LS}$  can be obtained by solving the equations

$$\left. \frac{\partial J(\theta_1, \theta_2)}{\partial \theta_1} \right|_{\theta_1 = \hat{\theta}_{1LS}} = 0 \quad \text{and} \quad \left. \frac{\partial J(\theta_1, \theta_2)}{\partial \theta_2} \right|_{\theta_2 = \hat{\theta}_{2LS}} = 0$$

❖ Thus,

$$\hat{\theta}_{1LS} = \bar{X} - \bar{a} \hat{\theta}_{2LS} \quad \text{and} \quad \hat{\theta}_{2LS} = \frac{\sum_{i=1}^N (X_i - \bar{X})(a_i - \bar{a})}{\sum_{i=1}^N (a_i - \bar{a})^2}$$

where

$$\bar{X} = \frac{1}{N} \sum_{i=1}^N X_i$$

$$\bar{a} = \frac{1}{N} \sum_{i=1}^N a_i$$

We will consider another example suppose the observations  $X_i$   $i$  going from 1 to  $N$  are related by this relationship  $X_i = \theta_1 + a_i \theta_2 + e_i$ . So this is for  $i = 1, 2$  upto  $N$  where  $\theta_1$  and  $\theta_2$  are unknown parameters and  $e_i$  is the observation noise then the  $\theta_1$ ,  $\theta_2$  is given by this summation  $X_i$  this is the data – this model  $\theta_1 - a_i \theta_2$  whole square  $i$  going from 1 to  $N$ .

So LS estimators two terms we have to estimate  $\theta_1$  hat LS and  $\theta_2$  hat LS and they are obtained by taking the partial derivative with respect to  $\theta_1$  that is  $= 0$  and taking the partial derivative with respect to  $\theta_2$  and that  $= 0$  and ultimately we can take the partial derivative and solve these two equations. We can get  $\theta_1$  hat LS  $= \bar{X} - \bar{a}$  into  $\theta_2$  hat LS and  $\theta_2$  hat LS  $= \frac{\sum_{i=1}^N (X_i - \bar{X})(a_i - \bar{a})}{\sum_{i=1}^N (a_i - \bar{a})^2}$ .

So these are the expression for  $\theta_1$  hat LS and  $\theta_2$  hat LS where  $\bar{X}$  and  $\bar{a}$  are the sample mean. So that way  $\bar{X} = \frac{\sum_{i=1}^N X_i}{N}$  and  $\bar{a} = \frac{\sum_{i=1}^N a_i}{N}$ .

**(Refer Slide Time: 14:23)**

### Matrix formulation of the linear model

- ❖ The general linear model for observed samples  $X_i, i = 1, 2, \dots, N$  can be written in the matrix form

$$X = A\theta + e$$

where  $\theta = [\theta_1, \theta_2, \dots, \theta_K]^T$ ,

$A = [a_{ij}]$  is an  $N \times K$  matrix with  $N > K$  and

$e = [e_1, e_2, \dots, e_N]^T$  is the zero-mean random observation error vector. The random errors are assumed to be uncorrelated.

- ❖ Assume  $A$  to be a full-rank matrix. In other words, the columns of  $A$  are linearly independent.

- ❖ The SSE is given by

$$\begin{aligned} J(\theta) &= \sum_{i=1}^N e_i^2 \\ &= \sum_{i=1}^N \left( X_i - \sum_{j=1}^K a_{ij} \theta_j \right)^2 \\ &= (X - A\theta)^T (X - A\theta) \\ &= X^T X - X^T A \theta - \theta^T A^T X + \theta^T A^T A \theta \end{aligned}$$

Now we will have a matrix formulation of the linear model. The general linear model of observed samples  $X_i$   $i$  going from 1 to  $N$  can be written in the matrix form as  $X = A$  matrix times  $\theta$  vector +  $e$  vector. This is the linear model in the matrix form. Here  $\theta$  equal to the parameter vector comprising of  $\theta_1, \theta_2$  upto  $\theta_K$ . So this is the  $\theta$  vector similarly  $A$  matrix is the matrix of  $a_{ij}$  it is a  $N/K$  matrix  $N$  is the number of observation.

And  $K$  is the number of parameters and with the assuming that  $N > K$ .  $e$  is again corresponding to each observation we have one error therefore  $e$  is an  $N$  dimensional error vector comprising of  $e_1, e_2$  upto  $e_N$   $e_i$  are zero mean random variable. Further, the random errors are assumed to be uncorrelated. They are uncorrelated that is covariance of any two element will be  $= 0$  and we also assume that they are of constant variance.

We assume  $A$  to be a full rank matrix. In other words, the columns of  $A$  are linearly independent. This is important for having the solution. The sum square error is given by  $J(\theta)$  that  $= e_i^2$   $i$  going from 1 to  $N$  and that  $= \sum_{i=1}^N (X_i - \sum_{j=1}^K a_{ij} \theta_j)^2$   $i$  going from 1 to  $N$  this is the cost function and this in matrix notation we can write  $X - A\theta$  transpose into  $X - A\theta$ .

So this summation we can write in the matrix form like this  $X - A\theta$  transpose into  $X - A\theta$  and if we expand this product we will get this as  $X^T X$  because transpose is here  $X$  transpose into  $X - X$  transpose into  $A\theta$  now transpose of  $A\theta$  is  $\theta^T A^T$   $A$  transpose. So that way  $\theta^T A^T A \theta$  into  $X +$  this - - will be  $+ \theta^T A^T X$   $A A$  transpose into  $\theta$ . So product of this transpose into  $A\theta$ .

(Refer Slide Time: 17:36)

**Normal equations and solution**

- ❖ Thus  $\hat{\theta}_{LS}$  is given by  $\hat{\theta}_{LS} = \arg \min_{\theta} J(\theta)$

$$\therefore \frac{\partial J(\theta)}{\partial \theta} \bigg|_{\theta = \hat{\theta}_{LS}} = 0$$
$$\Rightarrow -A^T X - A^T X + 2A^T A \theta = 0$$
$$\therefore A^T A \hat{\theta}_{LS} = A^T X$$

- ❖ The above equation is known as the normal equation.
- ❖ When  $\det(A^T A) \neq 0$ , the optimal solution is given by
$$\hat{\theta}_{LS} = (A^T A)^{-1} A^T X = A^+ X$$
- ❖ The matrix inversion makes LS estimators computationally complex.
- ❖ Recursive least squares method is a solution to this problem.

Now the estimator is given by  $\hat{\theta}_{LS} = \arg \min_{\theta} J(\theta)$  and we can obtain by the relationship that is  $\frac{\partial J(\theta)}{\partial \theta} \bigg|_{\theta = \hat{\theta}_{LS}} = 0$ . Now this is a partial derivative with respect to  $\theta$  vector that way we have to do the partial differentiation with respect to vector. So this we can carry out for this expression this is a constant term so this term will have zero partial derivative similarly corresponding to this term we will have  $A^T X$ .

Similarly, for second term then this term also we will have the derivative  $A^T X$  like that and the third term will be 2 times  $A^T A \theta$  and that must be  $= 0$ . So this is the result we get after partial derivative with respect to  $\theta$  vector. From this we get an equation because here  $A^T X$  again  $A^T X$  so it will be  $2 A^T X$  and there is  $2 A^T A \theta$ .

So 2, 2 will get cancelled therefore we will get this relationship.  $A^T A \hat{\theta}_{LS} = A^T X$ . This is the equation governing LS estimation. So that means we have to pre-multiply  $\hat{\theta}_{LS} / A^T A$  and right hand side is  $A^T X$  this equation is known as the normal equation matrix form of normal equation. When determinant of  $A^T A$  is not  $= 0$  the optimal solution is given by  $\hat{\theta}_{LS} = (A^T A)^{-1} A^T X$  and this we denote by  $A^+$  into  $X$ .

So this is the least square solution and here this matrix inversion is there matrix inversion makes LS estimators computationally complex because when this matrix dimension is very high inversion is a big problem. Recursive least squares method is a solution to this problem.

which we will be discussing in the next lecture.

(Refer Slide Time: 20:46)

**Pseudo inverse**

- ❖ The matrix  $A^+ = (A'A)^{-1}A'$  is known as the pseudo inverse and has importance in solution of linear equations.
- ❖ Suppose  $A = \begin{bmatrix} 1 & 0 \\ 0 & 1 \\ 0 & 2 \end{bmatrix}$ . Then,

$$A^+ = \left( \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 2 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & 1 \\ 0 & 2 \end{bmatrix} \right)^{-1} \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 2 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & \frac{1}{5} & \frac{2}{5} \end{bmatrix}$$

- ❖ When  $A$  is non-singular,  $A^+ = A^{-1}$

$$A^+ = (A'A)^{-1}A' = A^{-1}(A'A)A' = A^{-1}A = I$$

The quantity  $A^+$  that  $= A^T A^{-1} A^T$  is known as the pseudo inverse and has importance in solution of linear equation. In many cases of linear equations solution can be obtained in terms of pseudo inverse. Suppose  $A =$  this matrix, matrix comprising of 1 0 0 1 0 2. So 2 columns are there this 2 columns are independent and here pseudo inverse  $A^+$  that is  $A^+$  will be  $= A^T A^{-1} A^T$  so this is the expression for pseudo inverse.

So that way if I take the transpose of this I will get 1 0 0 0 1 2 and then  $A$  is this whole inverse multiplied by  $A^T A^T$  is 1 0 0 0 1 2. So if I carry out all this operations then we will get that pseudo inverse of  $A = 1 \ 0 \ 0 \ 0 \ 1/5 \ 2/5$ . When  $A$  is non singular pseudo inverse of  $A$  will be simply  $A^{-1}$  itself because in that case  $A^+$  will be  $= A^T A^{-1} A^T$ .

And in that case now we can inverse like this so we can write this as  $A^{-1}$  then  $A^T A^{-1} A^T$  so this is identify matrix so we will have simply  $A^{-1}$ . So when  $A$  is non singular so pseudo inverse is  $A^{-1}$  itself.

(Refer Slide Time: 23:25)

### Example 2 revisited

- The observations  $X_i, i = 1, 2, \dots, N$  are related by  $X_i = \theta_1 + i\theta_2 + e_i$

$$X = \begin{bmatrix} X_1 \\ X_2 \\ \vdots \\ X_N \end{bmatrix}, \quad A = \begin{bmatrix} 1 & 1 \\ 1 & 2 \\ \vdots & \vdots \\ 1 & N \end{bmatrix}$$

$$\therefore A'A = \begin{bmatrix} N & \sum_{i=1}^N i \\ \sum_{i=1}^N i & \sum_{i=1}^N i^2 \end{bmatrix}$$

- The normal equations are given by

$$A'A\hat{\theta}_{LS} = A'X$$

$$\therefore \begin{bmatrix} N & \frac{N(N+1)}{2} \\ \frac{N(N+1)}{2} & \frac{N(N+1)(2N+1)}{6} \end{bmatrix} \begin{bmatrix} \hat{\theta}_{1LS} \\ \hat{\theta}_{2LS} \end{bmatrix} = \begin{bmatrix} 1 & 1 \\ 1 & 2 \\ \vdots & \vdots \\ 1 & N \end{bmatrix} \begin{bmatrix} X_1 \\ X_2 \\ \vdots \\ X_N \end{bmatrix} = \begin{bmatrix} \sum_{i=1}^N X_i \\ \sum_{i=1}^N iX_i \end{bmatrix}$$

Solving the above matrix equation, we can get  $\hat{\theta}_{1LS}$  and  $\hat{\theta}_{2LS}$

Let us consider example 2 again  $X_i$  are related by  $X_i = \theta_1 + i\theta_2$  suppose  $e_i = i + e_i$ . So in that case  $X =$  this  $X$  factor is given by this  $A$  matrix is given by this and therefore  $A$  transpose  $A$  will get like this. First element will be  $N$  second element will be summation  $i$   $i$  going from 1 to  $N$  and this element will be = summation  $i$   $i$  going from  $i$  to  $N$  and this element will be summation  $i$  square  $i$  going from 1 to  $N$ .

This is the  $A$  transpose  $A$  so  $A$  is like this  $A$  transpose  $A$  will be like this. Therefore, we will have the normal equation now, normal equation is  $A$  transpose  $A$  multiplied by  $\hat{\theta}_{LS} = A$  transpose  $X$ . This is the normal equation for least square estimation and we know now  $A$  transpose  $A$  this is given by this matrix into  $\hat{\theta}_{LS}$  2 component  $\theta_1$  hat  $LS$ ,  $\theta_2$  hat  $LS$ .

And that must be =  $A$  transpose that is transpose of this matrix multiplied by  $X_1, X_2$  upto  $X_N$  this vector. So if we carry out this multiplication transpose of this matrix into this vector will get this right hand side will be = summation  $X_i$   $i$  going from 1 to  $N$  and here it will be summation  $i$   $X_i$   $i$  going from 1 to  $N$  because this column has 1, 2 upto  $N$  therefore if I multiply transpose of this by this I will get this expression.

So that way we have now this set of expression it is  $(( ))$  (25:34) matrix. Solving the above matrix equation, we can get  $\theta_1$  hat  $LS$  and  $\theta_2$  hat  $LS$ .

**(Refer Slide Time: 25:47)**

### Statistical properties of $\hat{\theta}_{LS}$

(1)  $\hat{\theta}_{LS}$  is unbiased:

We have

$$\begin{aligned}\hat{\theta}_{LS} &= (A'A)^{-1} A'X \\ &= (A'A)^{-1} A'(A\theta + e) \\ \therefore E\hat{\theta}_{LS} &= (A'A)^{-1} A'A(E\theta + Ee) \\ &= \theta\end{aligned}$$

$\therefore \hat{\theta}_{LS}$  is unbiased.

We will examine if  $\hat{\theta}_{LS}$  satisfies the minimum variance property.

Let us see this statistical properties of theta hat LS. Number one is theta hat LS is unbiased. We have theta hat LS that is by definition  $A$  transpose  $A$  whole inverse into  $A$  transpose  $X$  that =  $A$  transpose whole inverse into  $A$  transpose now we know that  $X = A\theta + \text{error}$  so that way we will have  $E$  of theta hat LS. Therefore, will be =  $A$  transpose  $A$  inverse  $A$  transpose into  $A$  into  $E$  of theta expected value of theta +  $E$  of  $e$ .

This part is 0 and this theta is a constant here so if I carry out this and this, this will be identity matrix and therefore we will get simply theta. So  $E$  of theta hat LS will be = theta. Therefore, theta hat LS is unbiased. So one important property that least square estimator is an unbiased estimator. So this we proofed here. We will examine if theta hat LS satisfies the minimum variance property that is important property we have to examine.

(Refer Slide Time: 27:26)

2) Theorem: The covariance matrix of  $\hat{\theta}_{LS}$  is given by

$$C_{\hat{\theta}_{LS}} = \sigma^2 A^+ A'^+$$

where  $\sigma^2$  is the variance of the measurement noise.

Proof: Note that  $A^+ A = AA^+ = I$

$$\begin{aligned}C_{\hat{\theta}_{LS}} &= E(\hat{\theta}_{LS} - \theta)(\hat{\theta}_{LS} - \theta)' \\ &= E(A^+ X - \theta)(A^+ X - \theta)' \\ &= E(A^+ (A\theta + e) - \theta)(A^+ (A\theta + e) - \theta)' \\ &= E(A^+ ee'A^+)' \\ &= A^+ \sigma^2 I A'^+ \\ C_{\hat{\theta}_{LS}} &= \sigma^2 A^+ A'^+\end{aligned}$$

$$\begin{aligned}A^+ A &= (A'A)^{-1} A'A \\ &= I \\ A'^+ A &= I = AA'^+\end{aligned}$$

❖ The diagonal elements of  $C_{\hat{\theta}_{LS}}$  give the variance of each component of  $\hat{\theta}_{LS}$

❖ To prove the minimum variance property, we consider another unbiased linear estimator  $\tilde{\theta} = DX$ .

We have proof that  $\hat{\theta}_{LS}$  is an unbiased estimator. We will prove the theorem the covariance matrix of  $\hat{\theta}_{LS}$  is given by  $C_{\hat{\theta}_{LS}} = \sigma^2 A^+ \text{ into } A^+$  transpose where  $\sigma^2$  is the variance of the measurement noise and  $A^+$  as you know it is this pseudo inverse of  $A$ . To prove this, first we note that  $A^+ A = A A^+ = I$ . So  $A^+ A$  that  $= A^T A^{-1} A^T$  that is  $A^+ \text{ into } A$ .

So we see that this part is the inverse of this part therefore this will be simply  $= I$  matrix identity matrix and we can prove the other way  $A^+ \text{ into } A$  will be also  $= I$ . This property is similar to the property of inverse that is  $A^{-1} \text{ into } A = I = A A^{-1}$ . So here this property is satisfied even by pseudo inverse. There  $A^+ A = A A^+ = I$ . Now let us see  $C_{\hat{\theta}_{LS}}$  by definition this is the expected value of  $\hat{\theta}_{LS} - \theta \text{ into } \hat{\theta}_{LS} - \theta$  transpose.

Now putting the value for  $\hat{\theta}_{LS}$  that  $= A^+ X$  therefore this expression will be  $= A^+ X - \theta \text{ into } A^+ X - \theta$  transpose. Now we know that  $X = A \theta + E$  error vector. Therefore, this expression will be  $= E$  of  $A^+ \text{ into } A \theta + e - \theta \text{ into } A^+ \text{ then } X = A \theta + e - \theta$  and whole transpose. So that way this expression can be written like this. Now we know that  $A^+ \text{ into } A = I$  similarly here also  $A^+ A = I$ .

And therefore this first part will be simply  $= \theta A^+ A \text{ into } \theta$  will be  $= \theta - \theta$  will get cancel. Similarly, here also  $A^+ A \text{ into } \theta$  itself this  $\theta$  and this  $\theta$  will get cancel. Therefore, this whole expression will be  $= E$  of  $A^+ e e^T \text{ into } A^+$  transpose because here transpose is there so it will come to the right and then with transpose. Now we know that this error vector is uncorrelated with constant variance  $\sigma^2$ .

Therefore, the expected value of  $e e^T$  will be  $= \sigma^2 \text{ into } I$ . Therefore,  $C_{\hat{\theta}_{LS}}$  will be  $= A^+ \sigma^2 I \text{ into } A^+$  transpose and this we can write as because  $\sigma^2 I$  if we multiply by  $I$   $I$  will get the same matrix so it will be  $\sigma^2 \text{ into } A^+ \text{ into } A^+$  transpose. So this is the covariance matrix of  $\hat{\theta}_{LS}$ . Now the diagonal elements of  $C_{\hat{\theta}_{LS}}$  give the variance of each component of  $\hat{\theta}_{LS}$ .

So that is important now this is a covariance matrix and the diagonal elements are the variance. Therefore, the diagonal elements of  $C_{\hat{\theta}_{LS}}$  gives the variance of each component of  $\hat{\theta}_{LS}$ . To prove the minimum variance property, we consider another

unbiased linear estimator  $\tilde{\theta} = DX$ . Suppose this is an unbiased estimator and with respect to this unbiased estimator we will discuss the minimum variance property.

**(Refer Slide Time: 33:03)**

**Theorem:** Suppose  $X = A\theta + e$  and  $\tilde{\theta} = DX$  is an unbiased linear estimator of  $\theta$ . Then  $DA = I$ .

Proof:

We have

$$E\tilde{\theta} = \theta$$

Also

$$E\tilde{\theta} = EDX = ED(A\theta + e)$$

$$= DA\theta$$

$$\therefore DA\theta = \theta$$

$$\Rightarrow DA = I$$

We will prove this result first. Suppose  $X = A\theta + \text{error vector}$  this is the model for LS estimation and  $\tilde{\theta} = DX$  is then unbiased linear estimator of  $\theta$  because it is a linear combination of  $X$  so therefore it is a linear estimator and it is assumed to be unbiased then  $D$  into  $A$  will be  $= I$  any linear estimator into  $D$  model matrix  $= I$ . We will prove the result we have the expected value of  $\tilde{\theta} = \theta$  because  $\tilde{\theta}$  is an unbiased estimator.

Also  $E$  of  $\tilde{\theta} = E$  of  $DX$  by definition  $\tilde{\theta} = DX$  and now put  $X = A\theta + e$  vector. Therefore,  $E$  of  $DX$  will be  $= E$  of  $D$  into  $A\theta + e$  vector and we know that  $E$  of  $e$  vector will be  $= 0$  vector. Therefore, we will be simply having  $DA$  into  $\theta$  because  $\theta$  is a constant quantity. So  $E$  of  $\tilde{\theta}$  will be  $= DA$  into  $\theta$ , but we know that  $E$  of  $\tilde{\theta} = \theta$ .

Therefore, this  $DA\theta$  must be  $= \theta$  which implies that this  $DA$  must be  $= I$  identity matrix. So that way we have proofed one important result that any linear unbiased estimator of  $\theta$  will satisfy this property  $DA = I$ .

**(Refer Slide Time: 35:09)**

**Theorem:**  $\hat{\theta}_{LS}$  is the best linear unbiased estimator (BLUE)

**Proof:**

Now,

$$\begin{aligned} C_{\tilde{\theta}} &= E(\tilde{\theta} - \theta)(\tilde{\theta} - \theta)' = E(DX - \theta)(DX - \theta)' \\ &= E(D(A\theta + e) - \theta)(D(A\theta + e) - \theta)' \\ &= EDe'e' = D E e e' D' \\ &= D \sigma^2 I D' \end{aligned}$$

$\tilde{\theta} = DX$

- ❖ The diagonal elements of  $C_{\tilde{\theta}}$  gives the variance of the linear estimators of individual parameters.
- ❖ We have to minimize the diagonal terms of  $DD'$  to find the best linear estimator.

Now we will proof our final result theta hat LS is the best linear unbiased estimator. So this is the estimator with a minimum variance among the class of linear unbiased estimator. We will explain this again proof  $C_{\tilde{\theta}}$  that covariance matrix of theta tilde = expected value of theta tilde – theta into theta tilde – theta transpose by definition and theta tilde = D so this we can write as E of theta tilde = D into X – theta into DX – theta transpose.

So we have written theta tilde = DX. Now we know that  $X = A\theta + e$  vector. So therefore we will get this expression = E of expected value of D into A theta + error vector – theta into D into A theta + error vector – theta transpose. Now we will be using the previous identity D into A = I. Therefore, this part will be D into A will be = I therefore this will be simply theta and this theta and this theta will get cancel.

And therefore we will be left with D multiplied by e here. Similarly, from this side also since D into A = I therefore this theta and – theta will get cancel and we will have D into e transpose. So  $(e e')$  (37:33) transpose will be given by e transpose into D transpose. So therefore this  $C_{\tilde{\theta}}$  of theta tilde = E of expected value of D into e e transpose into D transpose. Now we know that this e e transpose is a diagonal matrix we know expected value of e e transpose = sigma square times I matrix.

Therefore, here it will be simply D times sigma square I into D transpose and this will give me sigma square into D D transpose. Therefore, for any unbiased estimator theta tilde the covariance matrix  $C_{\tilde{\theta}}$  is given by sigma square into D into D transpose and the diagonal elements of  $C_{\tilde{\theta}}$  gives the variance of the linear estimators of individual

parameters. Suppose  $\text{DX} \tilde{\theta} = \text{DX}$  and this is a vector and its components will have some variance and those variances will be given by the diagonal elements of the covariance matrix.

We have to minimize the diagonal terms of  $\text{DD}^T$  to find the best linear estimator. We know that  $\text{D}$  is a linear estimator, but to be the best linear estimator the variance should be minimum therefore and we know the variance is given by the diagonal elements of  $\text{C} \tilde{\theta}$  therefore we have to minimize the diagonal terms of  $\text{DD}^T$  to find the best linear estimator.

(Refer Slide Time: 39:54)

**Theorem ...**

- ❖ Using the relation  $\text{DA} = \text{I}$  and some matrix algebra, we can show that

$$\text{DD}^T = \text{A}^T (\text{A}^T)^T + (\text{D} - \text{A}^T)(\text{D} - \text{A}^T)^T$$

- ❖ Thus,

$$\text{dia}(\text{DD}^T) = \text{dia}(\text{A}^T (\text{A}^T)^T) + \text{dia}((\text{D} - \text{A}^T)(\text{D} - \text{A}^T)^T)$$

*variance of the components of  $\tilde{\theta}$  = variance of components of LS estimator + non-negative terms*

- ❖  $\text{dia}(\text{DD}^T)$  is minimized, whenever those of  $\text{dia}((\text{D} - \text{A}^T)(\text{D} - \text{A}^T)^T)$  is minimized.
- ❖ If we put  $\text{D} = \text{A}^T$ , then  $\text{dia}((\text{D} - \text{A}^T)(\text{D} - \text{A}^T)^T) = 0$ .
- ❖ Therefore,  $\hat{\theta}_{LS} = \text{A}^T \text{X}$  has the lowest variance.

Now using the relation  $\text{DA} = \text{I}$  and some matrix algebra we can show that this is a very important relationship  $\text{DD}^T = \text{A}^T (\text{A}^T)^T + (\text{D} - \text{A}^T)(\text{D} - \text{A}^T)^T$ . So this is the covariance matrix of the least square estimators and this is some additional terms, but we know that this is a non-negative quantity. Thus, the diagonal elements of  $\text{DD}^T = \text{dia}(\text{A}^T (\text{A}^T)^T) + \text{dia}((\text{D} - \text{A}^T)(\text{D} - \text{A}^T)^T)$ . So this is the covariance matrix of the least square estimators and this is some additional terms, but we know that this is a non-negative quantity. Thus, the diagonal elements of  $\text{DD}^T = \text{dia}(\text{A}^T (\text{A}^T)^T) + \text{dia}((\text{D} - \text{A}^T)(\text{D} - \text{A}^T)^T)$ .

So what does it mean? Variance of the components of  $\tilde{\theta} = \text{variance of components of LS estimators} + \text{some non-negative terms}$ . Therefore, this term will be minimum whenever this part is minimum. Variance of components of  $\tilde{\theta}$  will be minimum when this part, the corresponding diagonal elements of  $\text{D} - \text{A}^T + \text{D} - \text{A}^T$ , is minimum. This is an important observation.

That is diagonal of  $D D^T$  is minimized whenever those of the diagonal elements of  $D - A^+$  are minimized. If we put  $D = A^+$  then the diagonal elements of  $D - A^+$  will be  $= 0$ . So this term will be lowest so when  $D = A^+$  that means when the least square estimator  $A^+ D = A^+$  then the variance will be minimum. Therefore, we conclude that  $\hat{\theta}_{LS} = A^+ X$  that is pseudo inverse of  $A$  into  $X$  has the lowest variance. Thus, we have proved that  $\hat{\theta}_{LS}$  has the lowest variance among the unbiased linear estimators.

(Refer Slide Time: 43:25)

### Discussion

- ❖ Thus  $\hat{\theta}_{LS} = A^+ X$  is the BLUE. It is an unbiased estimator with the lowest variance among the unbiased estimators in the form  $\tilde{\theta} = DX$ .
- ❖ Further, the BLUE becomes an MVUE if the data samples are jointly Gaussian. *MMSE is a linear combination of data vector.*
- ❖ The pseudo inverse  $A^+ = (A^T A)^{-1} A^T$  involves matrix inversion. This is the main drawback of the LS estimators when the dimension of  $A$  is large.
- ❖ The recursive LS method overcomes this drawback. We will discuss the RLS technique while discussing about the RLS adaptive filter.

Therefore,  $\hat{\theta}_{LS}$  that  $= A^+ X$  is the BLUE best linear unbiased estimator. It is an unbiased estimator with the lowest variance among the unbiased estimators in the form  $\tilde{\theta} = DX$ . We are considering estimators in this form only that is the linear estimators and with additional property that this estimator is unbiased. So this makes least square estimators very attractive.

It is the best linear unbiased estimator. Further, the BLUE becomes MVUE recall MVUE that is minimum variance unbiased estimator. So BLUE will become MVUE if the data samples are jointly Gaussian. Then MMSE minimum mean square error estimator is a linear combination of data vectors and here also this least square estimator is a linear combination of data vector.

Therefore, this estimator will become MMSE minimum mean square error estimator and further it is unbiased therefore it will be minimum variance unbiased estimator. A pseudo inverse  $A^+ = A^T A^{-1} A^T$  involves matrix inversion. So this is the

matrix inversion we have to do. This is the main drawback of the LS estimators when the dimension of  $A$  is large.

Because matrix inversion is computationally highly expensive. The recursive LS method overcome this drawback. There is a technique called recursive least squares estimation which overcomes this drawback. We will discuss the RLS technique while discussing about the RLS adaptive filter.

**(Refer Slide Time: 46:00)**

**Summary**

- ❖ In LS estimation, the observed data are represented as a known function of some unknown parameters plus some errors which are assumed to be random
- ❖ The general linear model considered for LS estimation is given by  

$$X = A\theta + e$$
- ❖  $A$  is assumed to be a full-rank matrix. In other words, the columns of  $A$  are linearly independent.
- ❖ The SSE is given by  

$$J(\theta) = (X - A\theta)'(X - A\theta) = X'X - X'A\theta - \theta'A'X + \theta'AA'\theta$$
- ❖ Minimizing the SSE gives the normal equation  

$$A'A\hat{\theta}_{LS} = A'X$$
- ❖ The LS estimator is given by  

$$\hat{\theta}_{LS} = (A'A)^{-1}A'X = A^+X$$
 where  $A^+$  is the pseudo inverse of  $A$ .

Let us summarize the lecture in LS estimation the observed data are represented as a known function some unknown parameters plus some errors which are assumed to be random. So only error is random otherwise it is a deterministic model. This general linear model considered for LS estimation is given by  $X = A\theta + e$  vector. So this is the model so  $\theta$  is the vector of unknown parameters  $A$  is a matrix and  $e$  is a error vector which is uncorrelated zero mean and with constant variance.

$A$  is assumed to be a full-rank matrix in other words the columns of  $A$  are linearly independent. This is required to have the least square solution. Now here the cost function is sum square error and it is given by  $J(\theta) = (X - A\theta)'(X - A\theta)$  and this we expanded as  $X'X - X'A\theta - \theta'A'X + \theta'AA'\theta$  this is the sum square error cost function.

Minimizing the SSE gives the normal equation  $A'A\hat{\theta}_{LS} = A'X$  this is the normal equation for LS estimation. Now the inverse of this matrix

exists therefore the LS estimator is given by  $\hat{\theta}_{LS} = (A^T A)^{-1} A^T X$  and this quantity is known as the pseudo inverse that we denote by  $A^+$  that is  $\hat{\theta}_{LS} = A^+ X$  where  $A^+$  is the pseudo inverse of  $A$ .

(Refer Slide Time: 48:35)

**Summary...**

- ❖  $\hat{\theta}_{LS}$  is an unbiased estimator  $E[\hat{\theta}_{LS}] = \theta$
- ❖ The covariance matrix of  $\hat{\theta}_{LS}$  is given by  $C_{\hat{\theta}_{LS}} = \sigma^2 A^+ A^{+T}$
- ❖ For any unbiased linear estimator  $\tilde{\theta} = DX$ ,  $DA = I$  and the covariance matrix  $C_{\tilde{\theta}} = \sigma^2 DD^T$
- ❖  $\hat{\theta}_{LS}$  is the best linear unbiased estimator (BLUE). Becomes an MVUE if data are Gaussian.
- ❖ LS estimation is computationally expensive because of the matrix inversion. This drawback is overcome by using the RLS algorithm.

With this background on LS estimation, we will discuss the RLS adaptive filters in the next class

We establish that  $\hat{\theta}_{LS}$  is an unbiased estimator. So  $E[\hat{\theta}_{LS}] = \theta$ . The covariance matrix of  $\hat{\theta}_{LS}$  is given by  $C_{\hat{\theta}_{LS}} = \sigma^2 A^+ A^{+T}$  where  $A^+$  is the pseudo inverse and  $\sigma^2$  is the variance of the error. For any unbiased linear estimator  $\tilde{\theta} = DX$  we establish that  $DA = I$   $I$  is the identity matrix and the covariance matrix  $C_{\tilde{\theta}}$  is given by  $\sigma^2 DD^T$ . This is also an important relation.

Next we establish using this expression we established that  $\hat{\theta}_{LS}$  is the best linear unbiased estimator BLUE. The BLUE becomes an MVUE if data are jointly Gaussian. So that we establish that if random data vectors are jointly Gaussian then least square estimator is the MVUE because it is a BLUE and when data are jointly Gaussian the BLUE will become MVUE. LS estimation is computationally expensive because of the matrix inversion.

Because determining the pseudo inverse involves matrix inversion and matrix inversion is computationally expensive. This drawback is overcome by using the RLS algorithm recursive least square algorithm that we will be explaining in the next lecture. With this background on the LS estimation we will discuss the RLS adaptive filter in the next lecture. Thank you.