

Introduction To Adaptive Signal Processing

Prof. Mrityunjay Chakraborty

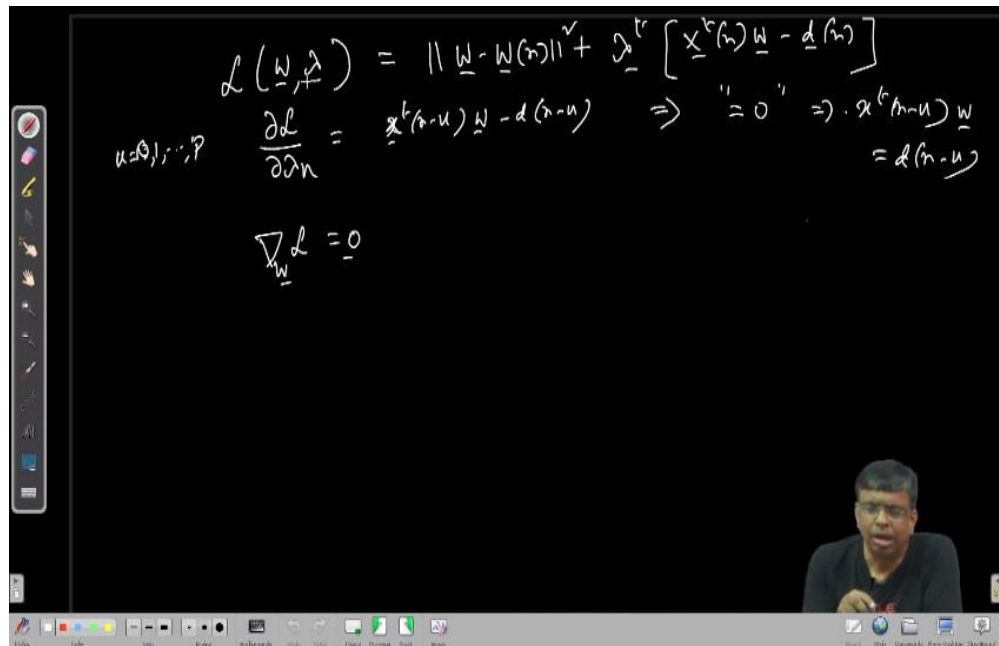
Department of Electronics and Electrical Communication Engineering

Indian Institute of Technology, Kharagpur

Lecture No # 31

Affine Projection Algorithm (APA)

So, in the last class we are carrying out the derivation of affine projection algorithm. So, this was the page we had a generalized version generalized function which was a function of all the weights w_0, w_1 up to w_{N-1} and also λ is a vector is a function of λ also λ_0, λ_1 up to λ_p because there are $p+1$ constraints ok. So, each constraint has one Lagrange multiplier and this was the thing all these were derived last time. So, I will not be wasting time in redoing this or restating this ok. This is a function we have to derive it with respect to each λ_k and we have to differentiate it with respect to each w ok.


$$\mathcal{L}(\underline{w}, \underline{\lambda}) = \|\underline{w} - \underline{w}(n)\|^2 + \sum_{k=0}^p \lambda_k [\underline{x}^T(n-k) \underline{w} - d(n-k)]$$
$$\frac{\partial \mathcal{L}}{\partial \lambda_n} = \underline{x}^T(n-n) \underline{w} - d(n-n) \Rightarrow "=0" \Rightarrow \underline{x}^T(n-n) \underline{w} = d(n-n)$$
$$\nabla_{\underline{w}} \mathcal{L} = \underline{0}$$

If I take any partial derivative of this with respect to λ_k λ_k figures here ok. So, actually it will be this you can see $x^T N d w$ ok. So, only this and λ transpose. So, this $\lambda_0 \dots \lambda_k \dots \lambda_p$ and this thing.

So, λ_k will multiply this $x^T N$ minus k vector. So, it is like this $x^T N$ minus k vector. So, ρ vector into w column vector minus dN this is k that will be multiplied by λ_0 , but there is no λ_k there. λ_k will multiply this term $x^T N$ minus k this ρ times w ok. That is there is a filter output at N minus k th time index with the filter coefficient vector w .

So, this output minus dN minus k alright this is the thing this will be multiplied by λ_k . So, when you take a partial derivative with respect to λ_k only this will result that is what we had obtained earlier ok and you equate that to 0. So, you get this. So, by this partial derivatives with respect to λ s I get this alright this equation. So, this will be satisfied.

So, therefore, because whatever be the w_N whatever is the w that must satisfy those p constants by forming this generalized function and introducing the Lagrange multiplier we see when you differentiate this L with respect to λ_k 's those conditions get satisfied right. That is if you want to find out minima of this function you have to take minima with respect to both λ s and w 's. So, there it shows that that minima will will automatically satisfy those constants these are the constant equations ok. For k equal to 0, k equal to 1, up to k equal to p . So, those will be satisfied.

$$L(\underline{w}, \lambda) = \|\underline{w} - \underline{w}(n)\|^2 + \lambda^T [\underline{x}^T(n) \underline{w} - \underline{d}(n)]$$

$$\frac{\partial L}{\partial w(n)} = \underline{x}^T(n-u) \underline{w} - \underline{d}(n-u) \Rightarrow " = 0 " \Rightarrow \underline{x}^T(n-u) \underline{w} = \underline{d}(n-u)$$

$$\nabla_{\underline{w}} L = 0$$

$$[\lambda_0 \dots \lambda_p]$$

$$\begin{bmatrix} \underline{x}^T(n) \\ \underline{x}^T(n-u) \end{bmatrix} \underline{w} - \begin{bmatrix} \underline{d}(n) \\ \underline{d}(n-u) \end{bmatrix}$$

Now comes the other derivative. Here w minus now L is if you take this w minus wN norm square we have seen in the context of LMS this is minus the two cross terms are same. So, twice w transpose wN plus norm I mean you can write either way either in terms of norm or just this w transpose w is norm square of w . Similarly, w transpose $N wN$ is norm square of wN I am writing like this you could write as norm square also. So, this is the part and then you have got that extra part right.

$$L(\underline{w}, \lambda) = \underline{w}^T \underline{w} - 2 \underline{w}^T \underline{w}(n) + \underline{w}^T(n) \underline{w}(n) + \lambda^T [\underline{X}^T(n) \underline{w} - \underline{d}(n)]$$

So, therefore, if I take this derivative this function w transpose w we have already seen in the context of N LMS algorithm. What is w transpose w ? w_0 square plus w_1 square plus dot dot w_k square plus dot dot like that. Any partial derivative with respect to w_k will be twice w_k . So, if you put all the partial derivatives in a stack twice w_0 twice w_1 twice w_2 dot dot this is 2 if you take common it is $2w$. So, this term will give rise to $2w$ here twice w transpose wN that means, $w_0 w_0 N$, $w_1 w_1 N$, dot dot dot dot.

So, if you differentiate with respect to the particular weight say w_k . So, it will be this twice

the term is twice $w_k w_k^T N$. So, w_k goes because of differentiation will be twice $w_k^T N$. So, if you now stack then this derivative. So, twice $w_0^T N$ twice $w_1^T N$ twice $w_2^T N$ dot dot dot which will be twice $w^T N$ this has no w .

So, after differentiation it goes here $\lambda^T \text{transpose } x \text{ transpose } w$ minus $\lambda^T \text{transpose } dN \text{ transpose } dN$ has no w . So, that will go away because of derivation $\lambda^T \text{transpose } x \text{ transpose } N w$ it is nothing, but $x^T N \lambda^T \text{transpose } w$. Because you know from basic matrix theory any matrix if you have A into B transpose is $B^T A^T$ transpose λ may be a vector, but it is also matrix is not it P plus 1 cross 1. So, $x^T N \lambda^T \text{transpose } w$ which is $\lambda^T \text{transpose } x \text{ transpose } N$ ok. So, that times w .

$$\lambda^T X^T(n) \underline{w} = (\underline{X}(n) \underline{\lambda})^T \underline{w}$$

So, if you call it if you give it any name say may be α . So, $\alpha^T w$ and now if you differentiate it with respect to each component of w and put those results in a stack you will get nothing, but α vector like we get $w^T N$ vector here all right α vector w^T transpose I am sorry I wrote wrongly here no it is ok this is fine this is fine. This term I am talking of twice $w^T \text{transpose } w^T N$ and that is same as twice $w^T \text{transpose } N w$ twice $w^T \text{transpose } N w$. So, $w^T \text{transpose } N w$ when differentiated with respect to w get $w^T N$ instead of $w^T N$ I have α . So, $\alpha^T w$ when differentiated with respect to w put in a stack all the derivatives it will be α like it was $w^T N$ there.

So, it will give rise to α , α is this which means this derivative will be $\lambda^T \text{transpose } \alpha$ sorry α is $x^T N \lambda$. So, it will be $x^T N \lambda$. Now, remember it is a vector equation not a scalar equation $x^T N$ is a matrix λ is a vector $\lambda_0 \lambda_1$ up to λ_p and this is also a gradient vector all right. This is the situation this I have to equate to 0 vector for minimization for minima all right. That means, if I take this to the right hand side what I have is twice this $x^T N \lambda$ λ vector ok.

Now, if I multiply pre multiply both sides by $x^T N$ that is twice 2 remains same. So, 2 is just kept in the front then $x^T N$ times this gradient vector here also $x^T N$ $x N$ lambda all right.

$$\nabla_w L = 2\underline{w} - 2\underline{w}(n) + \underline{X}(n)\underline{\lambda} = \underline{0}$$

$$2(\underline{w}(n) - \underline{w}) = \underline{X}(n)\underline{\lambda}$$

Now, this matrix will if this matrix is invertible $x^T N$ $x N$ it is a square matrix Hermitian matrix I will talk in detail about this. So, if this is invertible your lambda will be just you know inverse of this will come here twice inverse of this then $x^T N$ then this all right ok. So, let me go to the next page.

$$2\underline{X}^T(n)(\underline{w}(n) - \underline{w}) = [\underline{X}^T(n)\underline{X}(n)]\underline{\lambda}$$

The screenshot shows a video player with the title "Lecture-31 : Affine Projection Algorithm (APA)". The video content displays handwritten mathematical derivations on a blackboard background. The derivations include:

- Initial expression for the cost function $L(\underline{w}, \underline{\lambda})$ involving $\underline{w}$, $\underline{w}(n)$, $\underline{X}(n)$, and $\underline{\lambda}$.
- Calculation of the partial derivative $\frac{\partial L}{\partial \underline{w}}$ resulting in $\underline{X}^T(n)\underline{w} - d(n)$.
- Setting the gradient $\nabla_{\underline{w}} L = \underline{0}$.
- Rearranging terms to get $2(\underline{w}(n) - \underline{w}) = \underline{X}^T(n)\underline{\lambda}$.
- Final expression for the gradient: $\nabla_{\underline{w}} L = 2\underline{w} - 2\underline{w}(n) + \underline{X}(n)\underline{\lambda} = \underline{0}$.
- Matrix notation for the final result: $2(\underline{w}(n) - \underline{w}) = \underline{X}^T(n)\underline{\lambda}$ and $2(\underline{w}(n) - \underline{w}) = \underline{X}^T(n)\underline{\lambda}$.

The video player interface includes a progress bar at the bottom showing 8:54 / 32:22, and standard YouTube controls like play, volume, and share buttons.

Let us study this matrix $x^T N$ $x N$ what was our $x N$? It had a vector $x N$ vector then $x N$ minus 1 vector dot dot dot $x N$ minus how many p $x N$ was the filter input data vector at N th clock. So, it was x of N if size is N cross 1. So, this matrix is N cross p plus 1 ok. If

the columns are linearly independent that is there is no linear relation between them that is you cannot write any column as a linear combination of the others ok. Then we say it is full rank full column rank ok.

In that case $x^T N x N$ will be invertible how we have coming to that consider this matrix $x^T N x N$. So, this is Hermitian. So, obviously, if you call it A A^T is again if you take the transpose of it you get $x^T N$ this will get transpose come here and $x^T N$ will come here upper transpose. So, transpose cancels ok. So, A is A^T Hermitian not only that it is positive semi definite always why because if you take for any vector $c \neq 0$ $c^T A c$ it will be we can see that it will be non negative why because c^T you replace A by $x^T N x N$ into c it is a matrix times a column vector you call it d .

$$\underline{A} = \underline{X}^t(n)\underline{X}(n)$$

$$\underline{A}^t = \underline{X}^t(n)\underline{X}(n)$$

$$\underline{A} = \underline{A}^t$$

So, d is a column vector and what is this? This is $x N c^T$ this will be d^T is it not if you take d as $x N$ matrix into a c and c is a non-zero vector. So, $x N$ into c if you call it d vector if you get transpose of d it will be $c^T x^T$ that is what you have. So, d^T and d is norm square of d . So, norm square can never be negative. So, it is always non negative ok that is always true.

$$\underline{c} \neq 0$$

$$\|d\|^2 \geq 0$$

$$\underline{x}(n) = \begin{bmatrix} x(n) \\ x(n-1) \\ \vdots \\ x(n-N+1) \end{bmatrix}_{N \times 1}$$

$$\underline{X}(n) = \begin{bmatrix} x(n) & x(n-1) & \dots & x(n-p) \end{bmatrix}_{N \times (p+1)}$$

$$\underline{A} = \underline{X}^H(n) \underline{X}(n) : \text{Hermitian}$$

$$\underline{A}^H = \underline{X}^H(n) \underline{X}(n)$$

$$\underline{A} = \underline{A}^H$$

• Positive semidefinite always

For any vector $\underline{c} \neq 0$

$$\underline{c}^H \underline{A} \underline{c} = \underline{c}^H \underline{X}^H(n) \underline{X}(n) \underline{c}$$

$$= \underbrace{\underline{c}^H \underline{X}^H(n)}_{\underline{d}^H} \underbrace{\underline{X}(n) \underline{c}}_{\underline{d}}$$

$$= \|\underline{d}\|^2 \geq 0$$

But if it is not only greater than equal to 0 if it is strictly equal to 0 strictly greater than 0 then it is positive definite then we have seen in the very beginning of the lecture positive definite matrices have positive eigenvalues and they are invariable. Now when this equal to 0 occurs equal to 0 it will be 0 if now what is this norm of d square? Norm of d square is d1 square d2 square d what? xN into c ok. So, it will be N cross 1 vector xN is N cross p plus 1 you are multiplying by c ok c is p plus 1 cross N for any vector let me write. So, p plus 1 entries in a vector c you are multiplying by xN there are p plus 1 columns. So, resulting vector will be length N vector right that is why.

So, d is length N vector norm square of d is this if this is equal to 0 since every term is positive or 0 no negative no term is negative every term is either positive or 0 positive or 0 and they are adding there is no subtraction. So, they can only add they can only contribute positively. So, the whole thing has to be 0 every term has to be 0 d1 square has to be 0 d2 square it cannot happen that this is also positive this was a positive, but there is a minus sign here. So, they cancel that is not there it is all positive signs additions only. So, if the sum total is 0 and every term is non negative there is either 0 or greater than 0 the only possibility is everybody 0 right only then it is.

$$\underline{x}(n) = \begin{bmatrix} x(n) \\ x(n-1) \\ \vdots \\ x(n-N+1) \end{bmatrix}_{N \times 1}$$

$$\underline{X}(n) = \begin{bmatrix} x(n) & x(n-1) & \dots & x(n-P) \end{bmatrix}_{N \times (P+1)}$$

$$\underline{A} = \underline{X}^H(n) \underline{X}(n) : \text{Hermitian}$$

$$\underline{A}^H = \underline{X}^H(n) \underline{X}(n)$$

$$\underline{A} = \underline{A}^H$$

• Positive semidefinite always
(B.11.1)

For any vector $\underline{c} \neq 0$

$$\text{If } \|\underline{d}\|^2 = 0$$

$$\downarrow$$

$$d_1^2 + d_2^2 + \dots + d_N^2 = 0$$

$$\underline{d}_1 = \underline{d}_2 = \dots = \underline{d}_N = 0$$

$$\underline{c}^H \underline{A} \underline{c} = \underline{c}^H \underline{X}^H(n) \underline{X}(n) \underline{c}$$

$$= \underline{d}^H \underline{d}$$

$$\|\underline{d}\|^2 \geq 0$$

$$= 0$$

So, only if $\underline{d}$ vector is 0 vector only if $\underline{d}$ vector is a 0 vector then it will be equal to 0 it will be therefore, but $\underline{d}$ equal to 0 what does it mean $\underline{d}$ equal to 0 means $\underline{x}_N$ into $\underline{c}_0$ is $\underline{c}$ is 0 this $\underline{c}_0 \underline{c}_1$ let me write and $\underline{x}_N$ what does this matrix times column vector what does it mean it is basically linearly combining the columns by this coefficients this actually $\underline{c}_0$ times $\underline{x}_N$ vector $\underline{c}_1$ times $\underline{x}_N$ minus 1 vector like that. So, $\underline{c}_0$ times $\underline{x}_N$ vector $\underline{c}_1$ times dot dot dot $\underline{c}_P$ times plus guy ok that is all 0 this means there is a linear relation among the columns suppose $\underline{c}_0$ is non 0 ok. So, I keep this term $\underline{c}_0 \underline{x}_N$ on the left hand side take the others on the right hand side $\underline{c}_0$ non 0. So, divide both side by $\underline{c}_0$ that is possible because $\underline{c}_0$ is non 0 I am assuming. So, I will have $\underline{x}_N$ right hand side will be minus $\underline{c}_1$ by $\underline{c}_0 \underline{x}_N$ minus 1 minus $\underline{c}_2$ by $\underline{c}_0 \underline{x}_N$ minus 2 dot dot dot minus $\underline{c}_P$ by $\underline{c}_0 \underline{x}_N$ minus P .

So, $\underline{x}_N$ will be a linear combination of $\underline{x}_N$ minus 1 to $\underline{x}_N$ minus P . That means there is there is some linear relation amongst the columns we assume that no such linear relation exists we assume no such because it is all coming from random data there is no model inside. So, that the columns are always governed by an equation of that kind ok because the data is just randomly coming alright. I am giving $\underline{x}_N$ data random purely random I am not giving data where there is a hidden linear dependence relation like that amongst the

data samples amongst the data samples. So, that the columns they become linearly dependent ok.

If that be then obviously d cannot be 0 because the moment d is 0 it implies this which means the columns are linearly dependent that is there is a linear relation between them. That is I mean that something we can always say it would not occur because I am giving x_N data randomly there is no linear relation hidden linear relation amongst the data they are all coming independently you know because random purely randomly calculating. So, you can never have these columns generated by the data as linear dependence, but this may not happen if this number of row if number of row N becomes less than P plus 1 ok. Then columns are linearly dependent ok.

I will tell you why ok. That means there will be some linear relation between the columns and there you can by appropriately choosing c_0 c_1 to c_p you can find ok a linear combination which is d that is x_N times c d equal to 0 even if c is non-zero ok because of this linear relation that will always happen why? Because now every column vector is length capital N ok. So, every column vector lies in a space of dimension capital N . So, I can have maximum capital N number of axis which are called basis, basis vectors ok in an N dimensional space. So, you have more columns than the number of axis number of basis like suppose this is one axis this is one axis these two are fine equality x axis you call it y axis, but if you have another fellow here z in the same plane x y z ok that is dimension was 2 I had 2 vectors spanning this plane there if I pack one more fellow z then this set is linearly dependent because z I can always write as sometimes N times x something times y by the parallelogram law of vector addition. Similarly, every x_N is length capital N .

So, they are belonging to a space of capital N cross 1 vectors that has dimension capital N . So, it can have maximum capital N number of axis that is basis any extra vector will be writable as a linear combination of this basis that is of this axis capital N number of axis. So, there will be a linear dependence relation in that case. Therefore, we should always avoid we should always avoid avoid avoid this that is choose ok. Then we can always argue

that the data which from these columns you know is such data is generated randomly from outside and there is no linear relation then you know governing the data.

So, that there will be some linear relation between these columns and all that that is a reasonable assumption, but if N is less than p plus 1 you cannot make that assumption because then every column vector is an N dimensional space, but p total column vector is p plus 1 which is more than N . So, we need capital N number of axis that is basis vectors are taken out ok. There are still extra vectors left they can be written as a linear combination of those x basis vectors those x in vectors ok. So, this is important. So, now this is the equation this is λ for a minute this is λ .

Handwritten notes on a blackboard:

$$\underline{X}(n) = \begin{bmatrix} x(n) & x(n-1) & \dots & x(n-p) \end{bmatrix}_{N \times (p+1)}$$

$$\underline{X}(n) = \begin{bmatrix} x(n) \\ x(n-1) \\ \vdots \\ x(n-N+1) \end{bmatrix}_{N \times 1}$$

If $N < P+1 \rightarrow$ columns are linearly dependent.

$$\underline{A} = \underline{X}^T(n) \underline{X}(n) : \text{Hermitian}$$

$$\underline{A}^T = \underline{X}^T(n) \underline{X}(n)$$

$$\underline{A} = \underline{A}^T$$

Positive semidefinite always

For any vector $\underline{c} \neq 0$

$$\underline{c}^T \underline{A} \underline{c} = \underline{c}^T \underline{X}^T(n) \underline{X}(n) \underline{c}$$

$$= \underline{d}^T \underline{d}$$

$$\|\underline{d}\|^2 \geq 0$$

$$\underline{d} = 0$$

$$\underline{X}(n) \underline{c} = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{bmatrix}$$

$$c_0 x(n) + c_1 x(n-1) + \dots + c_p x(n-p) = 0$$

If $\|\underline{d}\|^2 = 0$

$$d_1^2 + d_2^2 + \dots + d_N^2 = 0$$

$$d_1 = d_2 = \dots = d_N = 0$$

So, now, I am assuming $\underline{X}^T \underline{X}$ is indeed positive definite and therefore, invertible. So, both left hand side and right-hand side I multiply by the inverse of this ok. So, I get λ vector $2 \times \text{transpose } N \times \text{transpose } \underline{X}^{-1} \times \text{transpose } N$ wN minus w ok. So, this is λ and the constraint that we had $\underline{X}^T \underline{X} \underline{w} = \underline{d}$ if I take you know all these partial derivatives and put in a stack we get we have discussed this at length earlier that this constraint will be satisfied it will be $\underline{X}^T \underline{X} \underline{w} = \underline{d}$.

So, I have got another thing that is $x^T N$ into w equal to dN all right $x^T N w$ is dN .

Then come to this equation $2 w^T N - w^T N \lambda$, λ I will put here $2 w^T N - w^T N \lambda$ ok.

$$2(\underline{w}(n) - \underline{w}) = \underline{X}(n)\underline{\lambda}$$

So, λ you put here. So, these two comes $x^T N$ this $x^T N x^T N$ inverse then this $x^T N w^T N - w$ all right.

$$= 2\underline{X}(n) \left(\underline{X}^t(n)\underline{X}(n) \right)^{-1} \underline{X}^t(n)(\underline{w}(n) - \underline{w})$$

So, from here we will find out we have to I mean apply this constraint in this I see the 2 cancels 2 and 2 cancels right.

So, $x^T N w$ equal dN . So, $x^T N w$ will be dN . So, what if you have that you will have 2 2 is already cancelled, I take a minus negative of both side 2 and 2 cancel then take negative of both sides. So, $w - w^T N$ remember this w will be by $w^T N$ plus 1. So, $w - w^T N$ is minus and minus sign I will observe here. So, do not need to write minus separately ok 2 2 cancels.

So, we have got $x^T N$ now $x^T N w^T N - w$ will come on that and $x^T N w$ is this dN . So, it will be $x^T N w$ because I have taken minus of both sides. So, that will be positive this is negative $x^T N w$ is dN . So, this is my dN vector minus $x^T N w^T N$ all right. So, if then w will be what there is w there is that that I will equate as $w^T N$ plus 1.

So, update equation will be this I am taking to the right-hand side plus $x^T N$ this vector dN is a vector $x^T N w^T N$ that is also a vector this vector let me give a name e_n . So, this is my e_n vector ok. Now what is e_n ? Now let us see e_n means dN minus this right. So, dN

was d_N d_N minus 1 dot dot dot dot d_N minus p and minus $x^T N$ was this was x vector then $x^T N$ minus 1 vector all these were done earlier $x^T N$ minus p vector into w this is my e_N .

$$\underline{y} = 2 (\underline{x}^T(n) \underline{x}(n))^{-1} \underline{x}^T(n) (\underline{w}(n) - \underline{w})$$

$$\underline{x}^T(n) \underline{w} = d(n) \quad \underline{w}(n) - \underline{w} = \underline{x}(n) \underline{y}$$

$$= 2 \underline{x}(n) (\underline{x}^T(n) \underline{x}(n))^{-1} \underline{x}^T(n) (\underline{w}(n) - \underline{w})$$

$$\underline{w} - \underline{w}(n) = \underline{x}(n) (\underline{x}^T(n) \underline{x}(n))^{-1} [d(n) - \underline{x}^T(n) \underline{w}(n)]$$

$$\underline{w}(n+1) = \underline{w}(n) + \underline{x}(n) (\underline{x}^T(n) \underline{x}(n))^{-1} e(n)$$

What does it mean that w instead of w now it is w_N . So, w_N is the filter coefficient vector at the N th clock right. Suppose I hold it at N th clock filter output will be $x^T N w_N$ I subtract d_N that from d_N I get the output error. But now suppose I use the same filter coefficient, but take a look back I say that suppose I go back use these coefficients only, but use it in the previous clock N minus 1 that time my input data vector was $x^T N$ minus 1. So, then filter output with these coefficients will be $x^T N$ minus 1 w_N . So, that if I subtract from d_N minus 1 that will be the next component of error.

You understand that filter coefficient is not changing at N th clock whatever filter coefficient I have suppose I am freezing and I am taking back I am looking back that suppose I now use these coefficients only not w_N minus 1, but this w_N , but still at the previous clock ok at the previous clock that time my filter input vector was $x^T N$ minus 1. So, filter output would be $x^T N$ minus 1 and this filter w_N and that I subtract from d_N minus 1. So, I get the corresponding error and same next is $x^T N$ minus 2. So,

if I use the same filter coefficients, but at go back and you know to N minus 2th clock then my filter input is x_{N-2} vector x^T_{N-2} w_{N-2} ok that is my filter output that I can subtract from d_{N-2} ok and this. So, this that is why it is called a posteriori error that is I have w_N and now I am going back to previous clocks using that filter coefficient recalculating the output error ok that better vector.

The image shows a handwritten derivation of the basic APA weight update formula on a blackboard. The equations are as follows:

$$\underline{\lambda} = 2 (\underline{x}^T(n) \underline{x}(n))^{-1} \underline{x}^T(n) (\underline{w}(n) - \underline{w})$$

$$\underline{x}^T(n) \underline{w} = d(n) -$$

$$2 (\underline{w}(n) - \underline{w}) = \underline{x}(n) \underline{\lambda}$$

$$= 2 \underline{x}(n) (\underline{x}^T(n) \underline{x}(n))^{-1}$$

$$\underline{w} - \underline{w}(n) = \underline{x}(n) (\underline{x}^T(n) \underline{x}(n))^{-1} [d(n) - \underline{x}^T(n) \underline{w}(n)]$$

$$\underline{w}(n+1) = \underline{w}(n) + \underline{x}(n) (\underline{x}^T(n) \underline{x}(n))^{-1} \underline{e}(n)$$

basic APA weight update

a posteriori error

$$\begin{bmatrix} d(n) \\ d(n-1) \\ \vdots \\ d(n-p) \end{bmatrix} = \begin{bmatrix} x(n) & x(n-1) & \dots & x(n-p) \\ x(n-1) & x(n-2) & \dots & x(n-p-1) \\ \vdots & \vdots & \ddots & \vdots \\ x(n-p+1) & x(n-p) & \dots & x(n-2p+1) \end{bmatrix} \underline{w}_{p \times 1}$$

So, this is the APA basic APA weight update formula. Some modifications will be done on this as before. So, that I will do in the next class. Thank you very much.