

Introduction To Adaptive Signal Processing

Prof. Mrityunjay Chakraborty

Department of Electronics and Electrical Communication Engineering

Indian Institute of Technology, Kharagpur

Lecture No # 24

Second Order Analysis of LMS Algorithm

So, we had seen earlier that the LMS algorithm the filter weights converge not absolutely, but only in mean to their respective optimal values as it has an index small n goes to infinity that is we need something called steady state that is for large value of n . Typically, we what we have there if you take say any k th weight we had already discussed this I am just repeating. So, $w_k(n)$ and you are going through iteration you are updating it via LMS algorithm ok, particular weight I am considering its optimal value may be here w_k^{opt} . So, initially it will fluctuate in any region you know if you go like this, but in the end, it will fluctuate in around the optimal one. So, if you measure the mean of this $w_k(n)$ at n very large n or as n goes to infinity like here the mean will be optimal weight.



So, we get convergence only in that sense then one can ask the question that how good it is because if it is still fluctuating over a wide range. So, at any iteration I get a value much

away quite different from $w_{k, opt}$ because it is fluctuating over a wide range how to control the range how to make sure that the range is small ok. So, for that we require a second order analysis that is, we will take the error between $w_k(n)$ and $w_{k, opt}$ that is called filter weight error this is for only k th component, all the components put together will be a weight error vector $v(n)$ we will take that and we will take the variance of each other ok. And sum all the variances that variance will give you the power of the fluctuation right because otherwise, it can I mean if you do not take variance if you take the just the value it can be highly positive highly negative when you add them is 0.

So, you can indicate that are on the average value is small. So, we are happy, but actually not because plus and minus both are there that is why they cancel each other, but the moment you take variance it is the difference between $w_k(n)$ and $w_{k, opt}$. So, that increment you are taking and you are taking the square of that. So, that positive increment you can get you get positive value positive power and their average. So, if that is less then the range of fluctuation will be less.

So, that is why the second order analysis that is variance analysis of the filter weight error ok that is what we have to we have to carry out and that is a very tedious exercise, but I will not go through all derivations because that is beyond the scope, but I will give some important results in between and then carry on with the remaining part of the derivation with that. So, that you can develop an insight all right. So, here I take the filter weight error vector not with respect to a particular component say k th component, but all together minus the corresponding opt all right $v(n) = w(n) - w_{opt}$ as you know we have got.

$$\underline{v}(n) = \underline{w}(n) - \underline{w}_{opt}$$

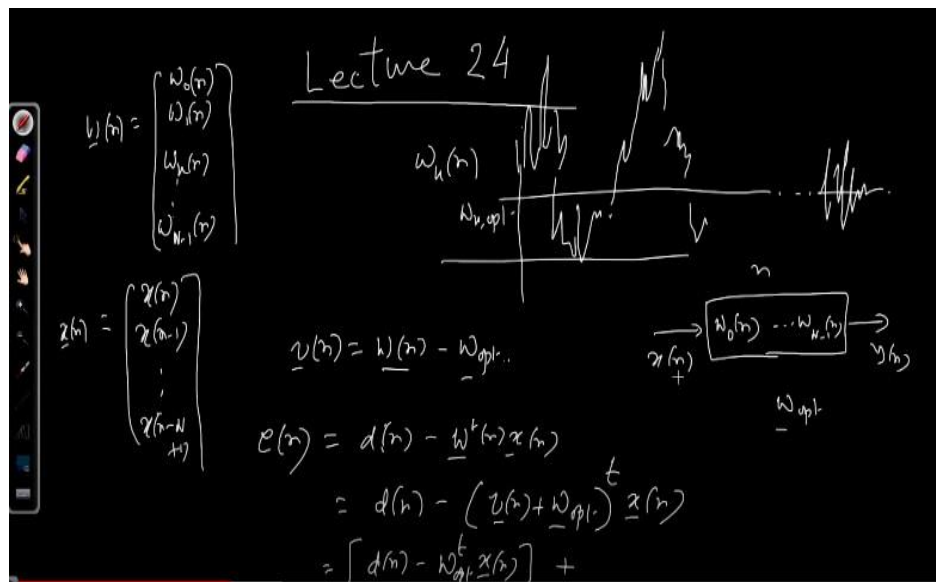
$$\underline{w}(n) = \begin{pmatrix} w_0(n) \\ w_1(n) \\ \vdots \\ w_{n-1}(n) \end{pmatrix}$$

So, all the weights you have w_k here and w_{opt} is a corresponding optimal vector. So, this will give you the filter weight error $w_k - w_{opt}$ filter error for k equal to 0, k equal to 1, k equal to k like that all right this component ok.

And in general, we are considering complex case. So, LMS algorithm in LMS or for that matter any adaptive filter says w_{n+1} or let me make it simpler. Suppose you have got an adaptive filter output error is e_n ok, not necessarily LMS what I am going to present is more general just an adaptive filter, filter coefficient vector at any n th clock is this you are updating them by some algorithm adaptive algorithm. So, at the n th clock the output error will be desired response minus as you have seen say $w^T x_n$. Let us consider real case because I remember I have not done possible in the complex LMS algorithm.

So, let us carry out with the real case. So, in the real case it will be $w^T x_n$ right, $w^T x_n$ this is a filter output, because this is a filter x_n goes in this you have seen many times filter output is y_n at any n th clock, these are the filter these are the filter coefficients. So, this $w^T x_n$ means and x_n is the data vector no change in definition I am just repeating them. So, $w^T x_n$ is $w_0 x_n + w_1 x_n + \dots$ which is the filter output and w_n is updated by some adaptive algorithm fine, but at any n th clock w_n is what if you take this w_{opt} to the left-hand side it is a w_{opt} plus the incremental part v_n . So, if you replace that here it will be $d_n - v_n + w_{opt}^T x_n$ alright, then you take this in one place what is $w_{opt}^T x_n$ that is if I have w_{opt} here then the corresponding filter output is $w_{opt}^T x_n$ ok.

Instead of this $w_0 - w_1 + \dots$ if I have w_{opt} coefficients like $w_0^{opt} w_1^{opt} w_k^{opt}$ and all that corresponding filter output will be that is $w_{opt}^T x_n$ like it is $w_{opt}^T x_n$ in general case x_n if w_n is w_{opt} it will be $w_{opt}^T x_n$.



So, when you are putting the optimal filter, you are going to the best case because y_n and d_n they are difference e_n has a minimum variance that is how w_{opt} was derived that you know optimal filter R inverse P that was derived by minimizing the variance of the output error between d_n and y_n ok. So, that is the best case that time y_n becomes a very good estimate of d_n because the error between them has been the minimum power minimum variance. So, that is this that is why I call it this also error output error filter output subtracted from d_n , but this is an optimal error e_{opt} let me call it e_{opt} and another term I have coming from the weight error vector term v_n transpose x_n alright. Therefore, what is the variance of this error variance.

$$E[e(n)] = E[e(n)e^t(n)] = E[(e_{opt}(n) - \underline{v}^t(n)\underline{x}(n))(e_{opt}(n) - \underline{v}^t(n)\underline{x}(n))^t]$$

So, this is the minimum variance because $e_{opt} n$ means output error variance is minimized then only you put $w_{opt} n$ and the corresponding error has minimum variance. So, if you calculate its variance, it will be a minimum thing you know I mean if you square up and take expectation dn square if you take expectation n will disappear because of w_{sn} . Similarly, cross terms also dn and xn elements of xn they are you know jointly stationary. So, again the cross correlations will have no n xn itself all the terms of xn they are $w_{ss} xn$ is w_{ss} process. So, any correlation between them and again independent of n that is why variance is just a constant independent of n minimum attainable.

$$\begin{aligned}
 e(n) &= e_{opt}(n) - \underline{v}^t(n)\underline{x}(n) & e_{opt}(n) &= d(n) - \underline{w}_{opt}^t \underline{x}(n) \\
 E[\tilde{e}(n)] &= E[e(n)e^t(n)] = E\left[(e_{opt}(n) - \underline{v}^t(n)\underline{x}(n))\right. \\
 &\quad \left.(e_{opt}(n) - \underline{v}^t(n)\underline{x}(n))^t\right] \\
 &= E[\tilde{e}_{opt}(n)]
 \end{aligned}$$

So, that is here what we will show is this except I mean in addition to the minimum attainable there will be additional components together will be this. This is happening because we are not using w_{opt} when you are calculating e_n we are using some arbitrary w_n which is updated by some adaptive algorithm. So, there e_n is not just $e_{opt} n$ and extra component coming due to the filter weight error that will contribute to some additional variance terms here that is what we evaluate this is fine. Then e_{opt} this cross term $\underline{v}^t(n)\underline{x}(n)e_{opt}(n)$ and minus actually $e_{opt}(n)\underline{v}^t(n)\underline{x}(n)$ and minus $e_{opt}(n)\underline{v}^t(n)\underline{x}(n)$ transpose $\underline{v}^t(n)\underline{x}(n)$ transpose $\underline{v}^t(n)\underline{x}(n)$ transpose. Now, $\underline{v}^t(n)\underline{x}(n)$ is a row vector this is a column vector.

So, you got a scalar and scalar transpose itself. So, this is actually equal to itself because this is a row vector this is a column vector row vector into column vector is a scalar any scalar and its transpose, they are same. So, this scalar transpose of the scalar is the scalar itself. So, you can replace this by v^T $n \times n$ v^T $n \times n$ is a scalar multiplied by another scalar e_{opt} n like here. So, these two terms are same and there is last term is expected value v^T $n \times n$ if I take transpose of this x^T v^T $n \times n$ v all right.

So, this is equal to minus these two terms are same. So, twice e let me write v^T $n \times n$ e_{opt} n plus all right. Now, to analyze further we make an assumption we make a more generalized we earlier obtained I had what is called independence assumption right. In the elimination analysis we assumed that the current filter weight vector w_n is statistically independent of the current data vector x_n and the reason we made the assumption that was also the motivation was clear that time I made it clear that time. We will just extend it here we will assume.

$$\begin{aligned}
 e(n) &= e_{opt}(n) - \underline{v}^T(n) \underline{x}(n) & e_{opt}(n) &= d(n) - \underline{w}_{opt}^T \underline{x}(n) \\
 E[\tilde{e}(n)] &= E[e(n) e^T(n)] = E\left[\left(e_{opt}(n) - \underline{v}^T(n) \underline{x}(n) \right) \left(e_{opt}(n) - \underline{v}^T(n) \underline{x}(n) \right)^T \right] \\
 &= \underbrace{E[\tilde{e}_{opt}(n)]}_{\tilde{\epsilon}_{min}} - E\left[\underline{v}^T(n) \underline{x}(n) e_{opt}(n) \right] \\
 &= \tilde{\epsilon}_{min} - 2 E\left[\underline{v}^T(n) \underline{x}(n) e_{opt}(n) \right] + E\left[\underline{v}^T(n) \underline{x}(n) \underline{x}^T(n) \underline{v}(n) \right] \\
 &= \tilde{\epsilon}_{min} - 2 E\left[\underline{v}^T(n) \underline{x}(n) e_{opt}(n) \right] + E\left[\underline{v}^T(n) \underline{x}(n) \underbrace{\underline{x}^T(n) \underline{x}(n)}_{= \underline{v}^T(n) \underline{x}(n)} \underline{v}(n) \right]
 \end{aligned}$$

We assume w_n is statistically independent of x_n of course, and also d_n all right. Logic is same if you really write w_n in terms of w_{n-1} w_{n-1} in terms of w_{n-2} and

all that and go up to w_0 which is an initial condition. I mean all the other terms they did not involve x_n and corresponding error e_n , but e_n will have again d_n and you know x_n I mean x_n all right and so on and so forth ok. So, this summation as I explained that time will be largely you know dominated by terms from the past because there is that is majority in number and past terms are usually uncorrelated with or even statistically independent with current terms that is why you can say w_n is statistically independent of this. If you do not understand you just go back to my lecture and that time go back to my lecture on LMS convergence analysis that time I made I explained what is meant by independence assumption.

So, it is just this ok just that time it was w_n was assumed to be independent of x_n ok. It was assumed to be independent of x_n and now d_n has been brought into it all right d_n has been brought into it. So, if that be now you see E of n it depends on d_n it depends on x_n w is constant here. That means, E of n is a function of d_n and x_n . So, here E of 10 is a function of x_n and d_n right consider v_n , v_n is nothing, but what is v_n ? v_n is nothing, but w_n minus w_{opt} .

So, basically you are subtracting a dc a constant from w_n ok. So, if w_n is statistically independent of x_n and d_n so is v_n , because v_n and w_n are equivalent it is just a dc shift of w_n by a constant w_{opt} . So, if w_n is statistically independent of x_n vector and d_n scalar so is v_n ok. So, here v_n it is statistically independent of x_n and E of n depends on d_n and x_n . So, it is a function of d_n and x_n this is the x_n itself.

So, this part is a function of x_n and d_n this is a vector x_n and this is a scalar E of n both are statistically independent with w_n and therefore, v_n E of n because E of n depends on d_n and x_n w is constant here and d_n and x_n they are statistically independent with w_n and therefore, v_n ok. So, E of n is statistically independent of v_n and therefore, $v^T n$ so is x_n by this assumption. So, x_n vector multiplied by the scalar E of n this resulting vector is statistically independent of this. So, expected value of so what does it mean? It means if I have suppose I need some space if I had suppose 2 vectors x and y they are S_i

statistically independent then E of x transpose y will be what E of x transpose y mean first component of x which is x_1 first component of y which is y_1 $x_1 y_1$ plus $x_2 y_2$ plus dot dot dot maybe $x_p y_p$ suppose length is 1 to p then E working on $x_1 y_1$ will be $E x_1 E y_1$ then $E x_2 E y_2 E x_p E y_p$ it is same as then becomes equivalent to E of x if you take E of x first that is E of x_1 , E of x_2 , E of x_3 this vector that transpose E of y , E of x_1 , E of y_1 . So, here you have got this vector E of x_1 here top guy is E of y_1 they are multiplied then next guy this is a row vector because of transposition E of x_2 here E of y_2 , E of x_2 , E of y_2 because statistically independent overall expectation of the product is expectation of x component into expectation of y component alright.

That means, this will be expectation of v transpose n into expectation of this ok, expectation of x_n into E of n this we can do. That means, it will be let me call it give a name this the variance this will be having this component minus twice have a look E of this and then E of this. You can do this transpose I think there was transpose. So, you can either apply E on v_n first then you take the transpose or first you do the transposition make it row vector then apply E on that either way it will be same into E of $x_n E$ of n plus this sorry E of v transpose $x x$ transpose v . Now, what is this term $x_n E$ of n we know this should be 0 the vector into scalar.

$$\begin{aligned}
 e(n) &= e_{opt}(n) - \underline{v}^t(n) \underline{x}(n) & e_{opt}(n) &= d(n) - w_{opt}^t \underline{x}(n) \\
 E[\tilde{e}(n)] &= E[e(n) e(n)] = E\left[\left(e_{opt}(n) - \underline{v}^t(n) \underline{x}(n) \right) \left(e_{opt}(n) - \underline{v}^t(n) \underline{x}(n) \right)^t \right] \\
 &= \underbrace{E[\tilde{e}_{opt}(n)]}_{\tilde{e}_{min}} - E\left[\underline{v}^t(n) \underline{x}(n) e_{opt}(n) \right] \\
 &= \tilde{e}_{min} - 2 E\left[\underline{v}^t(n) \underline{x}(n) e_{opt}(n) \right] + E\left[\underline{v}^t(n) \underline{x}(n) \underline{x}^t(n) \underline{v}(n) \right] \\
 &= \tilde{e}_{min} - 2 E\left[\underline{v}^t(n) \underline{x}(n) \underbrace{e_{opt}(n)}_{= \underline{v}^t(n) \underline{x}(n)} \right] + E\left[\underline{v}^t(n) \underline{x}(n) \underline{x}^t(n) \underline{v}(n) \right]
 \end{aligned}$$

Generalized Independence Assumption:

$$\underline{v}(n) \text{ is S.I. if } \underline{x}(n), d(n) \quad \left\{ \begin{array}{l} \underline{x}, y: \text{S.I.} \\ E\left[\underline{x}^t y \right] = E\left[\underline{x}^t \right] E\left[y \right] = E\left[\begin{matrix} x_1 y_1 + x_1 y_2 + \dots + x_1 y_p \end{matrix} \right] \\ = (E(\underline{x}))^t (E(y)) \end{array} \right.$$

$\underline{v}(n) = \underline{w}(n) - w_{opt}$

So, scalar times every component and this is 0 because under optimality condition we have seen the current optimal error it is orthogonal to there is uncorrelated to it each component of x_n ok. This again you can see you can write like you know $E[x_n]$ and E of n is d_n minus $W^T x_n$ and $W^T x_n$ is same as $x_n^T W$ and then x you write E of x_n into d_n that is p vector minus E of $x_n x_n^T W$ is constant. So, it goes outside E of $x_n x_n^T$ there is R and W . So, it is W opt under W opt because it is E of W opt and W opt is R inverse p .

So, p minus $R R$ inverse p . So, $R R$ inverse cancels p . So, it is 0 that this is exactly what we did last time all right. So, these two terms go. So, I am left with this much plus an extra component ok. So, because I did not use optimal weight optimal filter, I use just any arbitrary W_n .

So, overall output error variances gone up from the minimum attainable by a factor this extra factor. Remember one thing by your independence assumption W_n and therefore, V_n that is statistically independent of x_n right.

$$\begin{aligned} \epsilon_n^v &= E[\epsilon_n^v] = \epsilon_{min}^v - 2 E[\epsilon_{opt}(n)]^T E[x(n) \epsilon_{opt}(n)] \\ &\quad + E[x(n)^T x(n) x(n)^T \epsilon_{opt}(n)] \\ E[x(n) \epsilon_{opt}(n)] &= 0, \quad E[x(n) [d(n) - \underbrace{W_{opt}^T x(n)}_{x(n)^T W_{opt}}]] \\ &= 0 - (R W_{opt}) = p - R R^{-1} p = 0 \\ \epsilon_n^v &= \epsilon_{min}^v + E[x(n)^T x(n) x(n)^T \epsilon_{opt}(n)] \end{aligned}$$

Now, you can see one thing that suppose as an example a side story suppose I give you two vectors ok. A B is one vector may be α and I give another vector β as x y and it is

said alpha beta there is a sign suppose given this is V^T transpose sorry V transpose and it is given something like this E of like alpha transpose like V transpose alpha transpose then beta say beta transpose and then again alpha. You can do one thing beta beta I mean what will happen is this you have got $A B$ alpha transpose beta beta transpose and alpha.

So, you get $A B$ and beta beta transpose will be $x y$ column vector $x y$ row vector it will be $x^2 x y y x$ there is $x y y^2$ right. And now x^2 into A plus $x y$ into B multiplied by A on that if you apply expectation. So, that will be you know explain because alpha and beta are statistically independent E will get separated it will like you know I mean it will like this. This term E of $x^2 A$ plus $x y B$ into A plus again B into $x y A$ plus $y^2 B$ this E will acts separately on $x^2 x y$ again $x y$ and y^2 another E will be work on A into $A B$ and all that B into A it will be like that. You will see same thing if I keep one E like you know this is E with respect to both alpha beta.

Suppose I keep E with respect to this alpha out keep it as it is here you apply E beta x^2 into A ok. So, because E beta is only for $x y$ components. So, it will work out and I could separate out because they are statistically independent E alpha beta is E alpha E beta. So, E alpha beta if I apply on this, I take out E beta overall E over this will be what one expectation with respect to $x y$ that is beta another expectation with respect to alpha $A B$ they will separate out. So, the one with respect to beta I am directly applied here E beta $x^2 A$ again E beta $x y B$ and similarly E beta $x y A$ plus E beta $y^2 B$ alright this is what I am doing.

But this is same as though I am applying E beta here E beta E beta x^2 into A A into E beta x^2 into A again E beta $x y$ into B into A this is A . So, $A E$ beta $x^2 E$ beta x^2 into A plus E beta $x y$ into $B E$ beta $x y$ into B together times A plus B times B times within bracket E beta on $x y E$ beta on $x y$ into A plus E beta on y^2 into B . I am able to separate out E alpha beta as E alpha into E beta because they are statistically independent ok. So, then I can write like this I mean I can keep E alpha out and push like

this ok. So, on the inner thing this matrix I apply E beta which is a which is you know restricted to the elements of beta and then carry out the product then apply E alpha.

The same thing I can do here because here instead of alpha I have got VN. So, alpha transpose VN transpose and instead of beta I have got xN, but elements of VN and xN they are statistically independent because of the statistical independence assumption. So, same thing here so, which means I can write as E with respect to V if you want V transpose N E working on this and E working those will be input autocorrelation ok. E which is should be x working on this will be what nothing, but input autocorrelation. So, V transpose N or VN this product I have to now carry out which will be a scalar that have to take expectation with respect to VN elements of VN ok.

$$\begin{aligned} \epsilon_n^v &= E[\epsilon_n^v] = \epsilon_{min}^v - 2 E[v_{opt}^T(n)]^T E[x(n) e_{opt}(n)] \\ &\quad + E[v^T(n) x(n) x^T(n) v(n)] \\ E[x(n) e_{opt}(n)] &= 0, \quad E[x(n) [d(n) - \underbrace{w_{opt}^T x(n)}_{x^T(n) w_{opt}}]] \\ &= 0 - R w_{opt} = p - R R^{-1} p = 0 \\ \epsilon_n^v &= \epsilon_{min}^v + E[v^T(n) x(n) x^T(n) v(n)] \\ &= \epsilon_{min}^v + E_{v(n)} [v^T(n) R v(n)] \\ \frac{E_{x,R} [a(x^T a + x^T y) + b(y^T a + y^T y)]}{= E_x [R (E_R[x] a + E_R[y] b) + b (E_R[y] a + E_R[y] b)]} &= \begin{aligned} &\alpha = \begin{bmatrix} a \\ b \end{bmatrix} \quad \beta = \begin{bmatrix} x \\ y \end{bmatrix} \quad \alpha, \beta: \text{S.V.} \\ &E[\alpha^T \beta \beta^T \alpha] \\ &= E_x \left[\begin{bmatrix} a & b \end{bmatrix} E_{\beta} \begin{bmatrix} x^T & y^T \\ y & x \end{bmatrix} \begin{bmatrix} a \\ b \end{bmatrix} \right] \end{aligned} \end{aligned}$$

That will give me this extra component this I will analyze in the next class. Thank you very much.