

8.6 Convergence Behaviour of RLS

8.4.1

- To date, we have examined the principles and organization of RLS. Now it's time to check its convergence.

- Test bed for convergence study:

$$d(n) = e_o(n) + \underline{w}_0^T \underline{u}(n) \quad \text{and } \lambda = 1$$

where - measurement error $e_o(n)$ is zero mean, white, variance σ^2 , independent of $\underline{u}(i)$ for all i

- parameter vector $\underline{w}_0$ is constant

- where needed, $\underline{u}(i)$ is independent, identically distributed sequence (iid), with zero mean and cov matrix R

- Quantities of interest:

- mean value of estimate $\hat{\underline{w}}(n)$

- covariance matrix of estimate $\hat{\underline{w}}(n)$

- mean square value of $e(n)$, $J(n)$

- mean square value of $\hat{\underline{e}}(n)$, $J'(n)$, to compare with LMS.

- With startup by Solution 2 (block solution at time $n \geq M$) or approximately by Solution 3 (soft constrained, $\Phi^{-1}(0) = \delta^{-1} \mathbf{I}$), the RLS solution is the LS solution, just computed efficiently. Therefore these properties also apply to LS solutions generally, when applied to test bed.

• 8.6.1 The Solution Has the Right Mean

- We know from Section 7.4 of notes, Section 11.8 of Haykin, that LS estimates are unbiased if $e_o(i)$ process has zero mean. Therefore

$$E_e [\hat{w}(n)] = \underline{w}_0$$

- In order to obtain a useful relation en route, Haykin provides rederivation

$$\begin{aligned} \hat{w}(n) &= \Phi(n)^{-1} \underline{z}(n) = \Phi(n)^{-1} \sum_{i=1}^n \underline{u}(i) d^*(i) \\ &\quad \uparrow d(i) = e_o(i) + \underline{w}_0^T \underline{u}(i) \\ &= \Phi(n)^{-1} \sum_{i=1}^n \underline{u}(i) \underline{u}^T(i) \underline{w}_0 + \Phi(n)^{-1} \sum_{i=1}^n \underline{u}(i) e_o^*(i) \\ &= \underline{w}_0 + \Phi(n)^{-1} \sum_{i=1}^n \underline{u}(i) e_o^*(i) \end{aligned}$$

Since measurement error is assumed independent of $\underline{u}(i)$ process, and zero mean,

$$E[\hat{w}(n)] = \underline{w}_0 \quad \text{for } n \geq M \text{ (i.e., after block startup)}$$

- RLS (LS) is convergent in the mean for $n \geq M$, unlike LMS, where $E[\hat{w}(n)] \rightarrow \underline{w}_0$ for $n \rightarrow \infty$.

8.6.2 Covariance Matrix of the Solution $\underline{w}$

- We know from Section 7.4 of notes, Section 11.8 of Haykin, that covariance matrix of LS estimates is $\sigma^2 \Phi^{-1}(n)$.

- Haykin re-derives it as the start of a more useful conclusion. Define weight error vector

$$\underline{\epsilon}(n) = \underline{\hat{w}}(n) - \underline{w}_0 = \Phi^{-1}(n) \sum_{i=1}^n u(i) \mathbf{e}_0^*(i) \quad (\text{from 8.6.2})$$

- So cov matrix of the solution is

$$K(n) = E_{\underline{u}, \mathbf{e}_0} [(\underline{\hat{w}}(n) - \underline{w}_0)(\underline{\hat{w}}(n) - \underline{w}_0)^T] = E_{\underline{u}, \mathbf{e}_0} [\underline{\epsilon}(n) \underline{\epsilon}^T(n)]$$

$$= E_{\underline{u}} \left[\Phi^{-1}(n) \sum_{i=1}^n \sum_{j=1}^n u(i) \underbrace{\mathbf{e}_0^*(i) \mathbf{e}_0(j)}_{\sigma^2 \delta_{ij}} u^T(j) \Phi^{-1}(n) \right]$$

$$= \sigma^2 E_{\underline{u}} \left[\Phi^{-1}(n) \underbrace{\sum_{i=1}^n u(i) u^T(i)}_{\Phi(n)} \Phi^{-1}(n) \right] = \sigma^2 E_{\underline{u}} [\Phi^{-1}(n)]$$

- Use of $u(i)$ stats gives

$$K(n) = \frac{\sigma^2}{n-M-1} R^{-1} \quad n > M+1$$

(see Haykin 13.6)

assuming $u(i)$ iid, zero mean, cov R

- This makes the m.s. misadjustment

$$E[\underline{\epsilon}^T(n) \underline{\epsilon}(n)] = \text{tr}[K(n)] = \frac{\sigma^2}{n-M-1} \sum_{i=1}^M \frac{1}{\lambda_i}$$

(since eigenvals of R^{-1} are reciprocals of eigenvals λ_i of R)

Hence - decreases asymptotically as $1/n$

- small eigenvals λ_i of R appear to cause problems

- converges in mean square, no jitter in $\underline{\hat{w}}(n)$

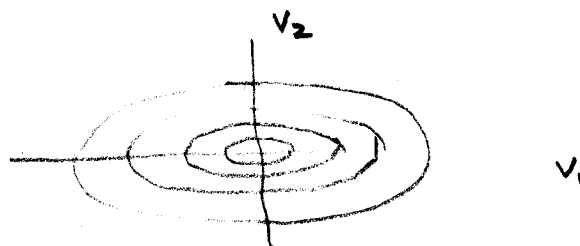
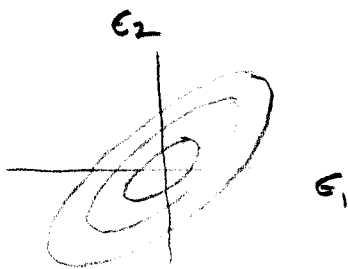
- excess mse?

— A clearer way to view this result is to recall that the small eigenvalue directions are less coupled to the output, so the error in estimating $d(n)$ may not be excessive.

— As in LMS analysis, expand using eigenvector basis:

$$\underline{\epsilon}(n) = \Phi \underline{v}(n) \quad \underline{v}(n) = \Phi^+ \underline{\epsilon}(n)$$

$$E[\underline{v}(n) \underline{v}(n)^+] = \frac{\sigma^2}{n-M-1} \Phi^+ R^{-1} \Phi = \frac{\sigma^2}{n-M-1} \begin{bmatrix} \lambda_1^{-1} & & 0 \\ & \lambda_2^{-1} & \\ 0 & & \ddots \\ & & & \lambda_M^{-1} \end{bmatrix}$$



The standard deviation along each axis is $\propto \frac{1}{\sqrt{\lambda_i}}$

This is exactly what we achieved in LMS after orthogonalization, so we expect RLS to be similarly insensitive to eigenvalue spreads

• 8.6.3 Mean Squared Estimation Error (Learning Curve)

— For many purposes, the real test of convergence is the estimation error. Use $\hat{\xi}(n)$, as in LMS analysis. Since we initialize with $\hat{\mathbf{w}}(0) = \mathbf{0}$, the ms error $E[|\hat{\xi}(1)|^2] = |d(1)|^2$ but should decrease from that value.

- The a priori estimation error at time n is

$$\begin{aligned}\xi(n) &= d(n) - \hat{w}^T(n-1) \underline{u}(n) \\ &= e_o(n) + \underline{w}_o^T \underline{u}(n) - \hat{w}^T(n-1) \underline{u}(n) \\ &= e_o(n) - \underline{\varepsilon}^T(n-1) \underline{u}(n)\end{aligned}$$

- The ms error is defined

$$\begin{aligned}J'(n) &= E[|\xi(n)|^2] \quad \text{although remember the cost function is } \sum_{i=1}^n |e(i)|^2 \\ &= \overline{|e_o(n)|^2} - \overline{\underline{\varepsilon}^T(n-1) \underline{u}(n) e_o^*(n)} - \overline{e_o(n) \underline{u}^T(n) \underline{\varepsilon}(n-1)} \\ &\quad + \overline{\underline{u}^T(n) \underline{\varepsilon}(n-1) \underline{\varepsilon}^T(n-1) \underline{u}(n)}\end{aligned}$$

Terms 2 and 3 are zero because $e_o(n)$ is indep of $\underline{u}(n)$ and all past values (test bed definition) and zero mean.

So

$$J'(n) = \sigma^2 + \overline{\underline{u}^T(n) \underline{\varepsilon}(n-1) \underline{\varepsilon}^T(n-1) \underline{u}(n)}$$

Time to invoke the independence assumption: $\underline{\varepsilon}(n-1)$ depends on $\underline{u}(i), d(i), i < n$ and $\underline{u}(n)$ is assumed iid (nonsense, of course, since $\underline{u}(n)_k = \underline{u}(n-1)_{k-1}, k=2, \dots, M$, but we'll run with it). Therefore $\underline{\varepsilon}(n-1)$ is indep of $\underline{u}(n)$.

$$J'(n) = \sigma^2 + E_{\underline{u}} \left[\underline{u}^T(n) E_{\varepsilon} [\underline{\varepsilon}(n-1) \underline{\varepsilon}^T(n-1)] \underline{u}(n) \right]$$

$$= \sigma^2 + \frac{\sigma^2}{n-M-1} \underbrace{E_{\underline{u}} [\underline{u}^T(n) R^{-1} \underline{u}(n)]}_{= \text{tr}[R^{-1}R]}$$

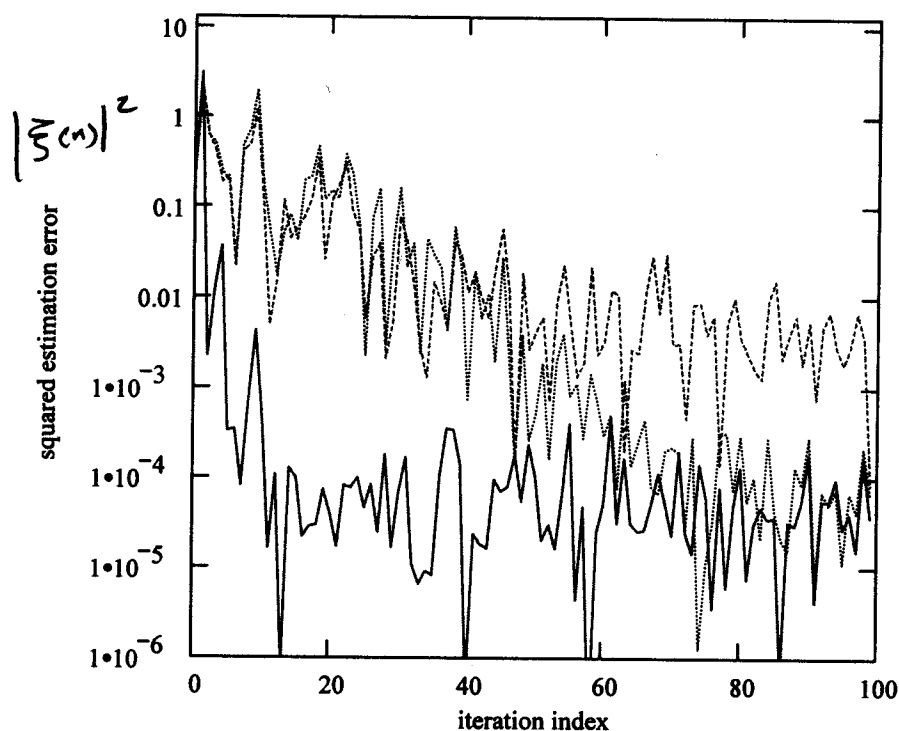
$$= \sigma^2 + \frac{\sigma^2 M}{n-M-1} = \sigma^2 \left(\frac{n-1}{n-1-M} \right) \quad n > M+1$$

No excess mse — $J'(n) \rightarrow \sigma^2 = J'_{\min}$

No dependence on eigenvalue spread!

• Summary:

- The ensemble mean of the solution vector is correct $E[\hat{w}(n)] = \underline{w}_0$ for all $n \geq M$
- The jitter in weights $K(n) = E[\underline{e}(n)\underline{e}^T(n)]$ is shaped, so that less jitter in directions strongly coupled to the estimation error. Unlike standard LMS, but like orthogonalized LMS
- The jitter in weights converges in mean square inversely with time to zero as $n \rightarrow \infty$. However, if $\lambda < 1$, then this isn't quite true (but we want some tracking ability).
- The a priori mse $J'(n)$ convergence is insensitive to eigenvalue spread
- By $n = 2M+1$, the ensemble averaged $J'(n)$ is only twice the minimum value $J'(2M+1) = 2\sigma^2 = 2J_{\min}$ from a start value of $J'(1) = E[\|d(1)\|^2]$. $J'(\infty) = J_{\min}$ for $\lambda = 1$.
- Comparison with LMS?
 - if small eigenval ratio, then about the same speed
 - if large spread, RLS unchanged, LMS much slower
 - RLS achieves lower error than LMS, since no excess mse
 - LMS can track for non zero μ ; RLS needs $\lambda < 1$ to track, so they should be compared on basis of same tracking ability



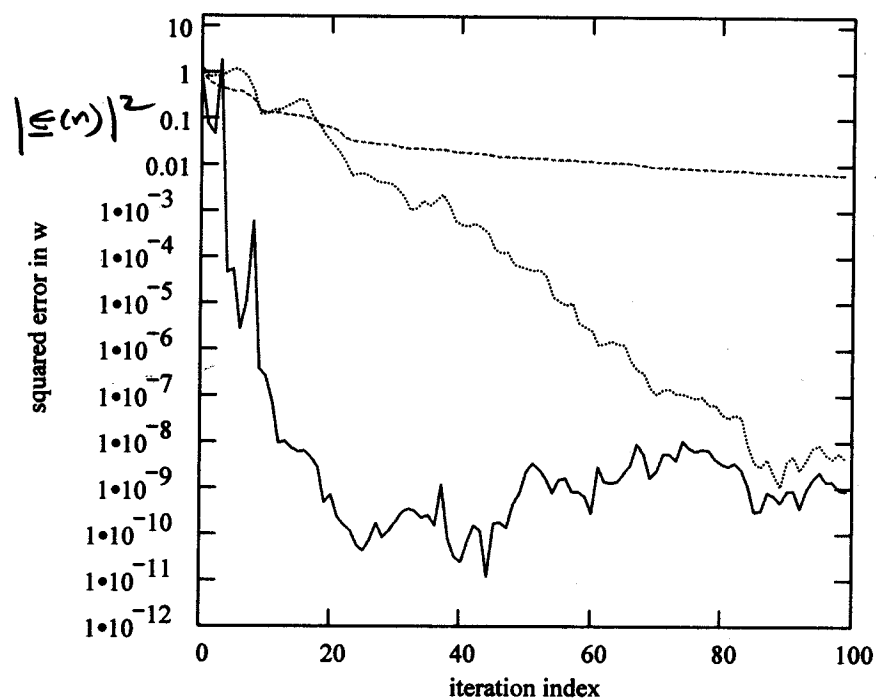
$$w = 0.45$$

$$M = 3$$

$$\frac{\max(\lambda)}{\min(\lambda)} = 31.4$$

$$\delta = 0.063$$

$$J_{\min} = 6.607 \cdot 10^{-5}$$



$$\lambda = 0.95$$

$$\gamma = 1 \cdot 10^{-3}$$

- Here's a summary of our set of methods for adaptation

		time	
		block	recursive
order	block	- estimate $r_u(n), r_{ud}(n)$; - solve $R \underline{w} = \underline{p}$; - mmse - LS; use $\underline{u}, \underline{d}$; solve $\Phi \underline{w} = \underline{z}$; - minimize $\ \underline{e}\ ^2$	- LMS; converges to mmse, but not a solution until convergence - RLS; a solution at every step; $\min \ \underline{e}\ ^2$ converges to mmse
	recursive	block adaptive lattice (Burg criterion); solution for all orders	- gradient lattice; converges to mmse; not a solution en route - RLS lattice; a solution at every step; converges to mmse; solution all orders