

6.4 Convergence and Excess Mean Squared Error

6.4.1

- The LMS algorithm, unlike steepest descent, has noisy inputs, so trajectory is a random walk with drift.



Because of shallow gradient near optimum, the gradient noise causes $w(n)$ to meander. The resulting misadjustment causes excess mse

$$J(n) = J_{\min} + \underbrace{\epsilon^T(n) R \epsilon(n)}_{J_{\text{ex}}}$$

$$J_{\text{ex}} = E[\epsilon^T R \epsilon]$$

- Haykin 9.4 provides a thorough analysis of convergence behaviour. We will take a shortcut to excess mse in class.

— JG Proakis and JH Miller, "An Adaptive Receiver for Digital Signaling Through Channels with Intersymbol Interference" *IEEE Trans Inf Th*, vol IT-15, no 4, pp 484-497, July 1969.
(adapted from Widrow's work in 1966)

- In order to obtain $J_{ex} = E[\underline{\epsilon}^T R \underline{\epsilon}]$, we need the statistics of $\underline{\epsilon}$; i.e. its covariance matrix $E[\underline{\epsilon} \underline{\epsilon}^T]$. We obtain this from the update equation, a ΔE with noisy input.

Recall

$$\underline{w}(n+1) = \underline{w}(n) + \underbrace{\mu \underline{e}^*(n) \underline{u}(n)}_{\text{represent as mean value plus zero mean noise}}$$

$$= (\mathbf{I} - \mu R) \underline{w}(n) + \mu \underline{p} + \mu \underline{\delta}(n) \quad \underline{p} = R \underline{w}(n) + \underline{\delta}(n)$$

Shift to misadjustment vector $\underline{\epsilon}(n) = \underline{w}(n) - \underline{w}_0$

$$\underline{\epsilon}(n+1) = (\mathbf{I} - \mu R) \underline{\epsilon}(n) + \mu \underline{\delta}(n)$$

This is progress, but now we need the statistics of the noise $\underline{\delta}(n)$; i.e., $E[\underline{\delta}(n) \underline{\delta}^T(n-k)]$

- To obtain this noise covariance matrix, we'll make approximations valid at or in the neighbourhood of the convergence point

$$\begin{aligned} \text{— first, } \underline{e}^*(n) \underline{u}(n) &= (\underline{p} - R \underline{w}(n)) + \underline{\delta}(n) \\ &\approx \underline{\delta}(n) \end{aligned}$$

= 0 at convergence point

so we evaluate the covariances in $E[\underline{\delta}(n) \underline{\delta}^T(n)]$ at the convergence point and assume that they don't change much in that region

$$E[\underline{\delta}(n) \underline{\delta}^T(n)] = E[|\underline{e}(n)|^2 \underline{u}(n) \underline{u}^T(n)]$$

- next, assume $\underline{e}(n)$ is independent of $\underline{u}(n)$ and $\underline{u}(n) \underline{u}^T(n)$; true if Gaussian by principle of orthogonality; approx true if not Gaussian, since at least uncorrelated; and $\underline{e}(n)$ is approx Gaussian anyway by central limit theorem: $\underline{e}(n) = \underline{d}(n) - \underline{w}^T \underline{u}$

Using this independence, we obtain

$$E[\underline{\delta}(n) \underline{\delta}^T(n)] \approx E[|e(n)|^2] E[\underline{u}(n) \underline{u}^T(n)]$$

$$= J_{\min} R$$

↑ near convergence point

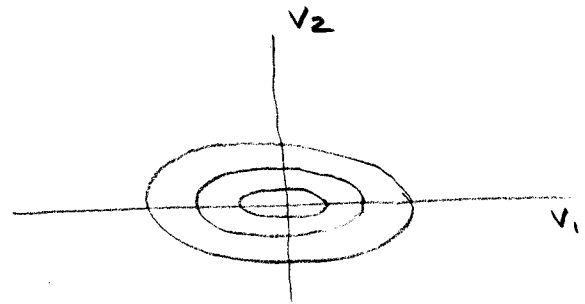
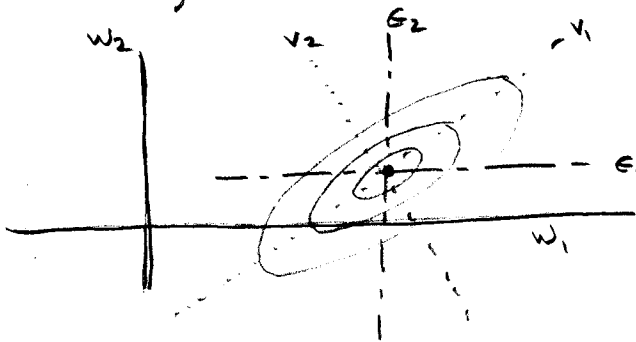
- and finally, assume that $\underline{\delta}(n)$ is a white sequence, because:
 - optimized estimation filters w exploit temporal correlations
 - if there were significant correlation between $\underline{\delta}(n)$ and $\underline{\delta}(n-k)$, the filter would adapt to use it, thereby removing it from $\underline{\delta}(n)$

- We now have a model for the noise driving the ΔE

$$\underline{\epsilon}(n+1) = (I - \mu R) \underline{\epsilon}(n) + \mu \underline{\delta}(n)$$

and we could obtain $R_e = E[\underline{\epsilon} \underline{\epsilon}^T]$ from it - but...

- it's easier and more instructive to transform to the eigenvector basis



$$\underline{\epsilon} = Q \underline{v} \quad Q = [q_1 | q_2 | \dots | q_M] \quad \underline{v} = \begin{bmatrix} v_1 \\ v_2 \\ \vdots \\ v_M \end{bmatrix}$$

v_i is the coefficient of mode i

- In terms of coefficients in the eigenvector basis

6.4.4

$$\underline{Q}(n+1) = (\mathbf{I} - \mu \mathbf{R}) \underline{Q}(n) + \mu \underline{\delta}(n)$$

$$\underline{v}(n+1) = (\mathbf{I} - \mu \mathbf{\Lambda}) \underline{v}(n) + \underbrace{\mu \mathbf{Q}^T \underline{\delta}(n)}_{\underline{z}(n)}$$

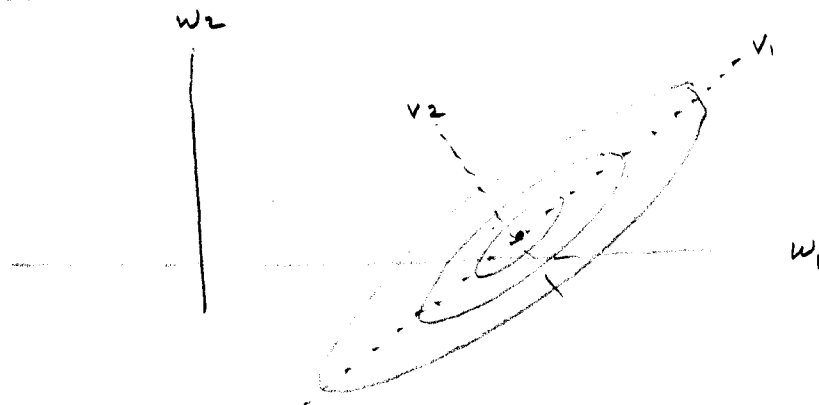
and the transformed noise covariance matrix is

$$\mathbf{E}[\underline{z} \underline{z}^T] = \mathbf{Q}^T \mathbf{E}[\underline{\delta} \underline{\delta}^T] \mathbf{Q} = \mathbf{J}_{\min} \underbrace{\mathbf{Q}^T \mathbf{R} \mathbf{Q}}_{\text{diagonal}} = \mathbf{J}_{\min} \mathbf{\Lambda}$$

so transformed noise components are uncorrelated,
and have variances $\sigma_{z_i}^2 = \mathbf{J}_{\min} \lambda_i$

- Transformation has expressed misadjustment dynamics as a set of uncoupled modes:

$$\boxed{v_i(n+1) = (1 - \mu \lambda_i) v_i(n) + \mu z_i(n)} \quad i=1, \dots, M$$



- Two observations:

- $(1 - \mu \lambda_i)$ is smallest for the dominant mode, the one with $\lambda_{\max}$, so it decays most quickly, less opportunity for noise buildup

- but $\sigma_{z_i}^2$ is proportional to λ_i , so greatest for $\lambda_{\max}$, because steepest in that direction, more $\text{en}(\underline{u}(n))$ variability.

- Which effect will win?

- Now we can obtain the coefficient variance in each mode from the transformed ΔE :

$$v_i(n+1) = (1 - \mu \lambda_i) v_i(n) + \mu v_i(n)$$

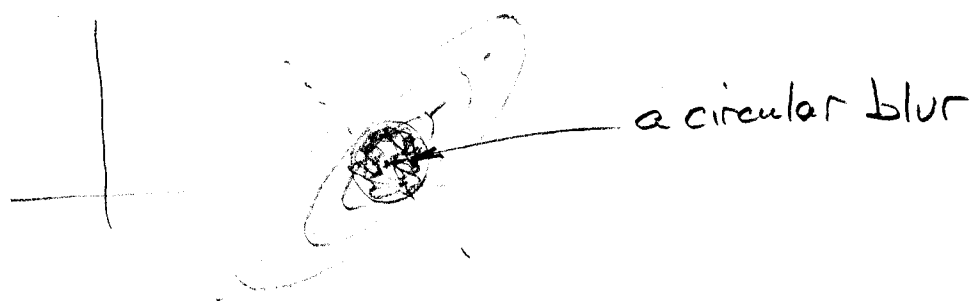
$$|v_i(n+1)|^2 = (1 - \mu \lambda_i)^2 |v_i(n)|^2 + \mu^2 |v_i(n)|^2 + 2(1 - \mu \lambda_i) \operatorname{Re}[v_i(n) v_i^*(n)]$$

$$\sigma_{v_i}^2 = (1 - \mu \lambda_i)^2 \sigma_{v_i}^2 + \mu^2 \sigma_{v_i}^2$$

$$\sigma_{v_i}^2 = \frac{\mu^2 \sigma_{v_i}^2}{1 - (1 - \mu \lambda_i)^2} = \frac{\mu^2 J_{\min} \lambda_i}{1 - (1 - \mu \lambda_i)^2}$$

$$\approx \frac{\mu J_{\min}}{2} \text{ for } \mu \lambda_i \ll 1$$

Amazing — same for all modes!



Therefore

$$E[\underline{v} \underline{v}^+] = \frac{\mu J_{\min}}{2} \mathbf{I}$$

and the quantity we were after all along...

$$E[\underline{\epsilon} \underline{\epsilon}^+] = \mathbf{Q} E[\underline{v} \underline{v}^+] \mathbf{Q}^+ = \frac{\mu J_{\min}}{2} \mathbf{I}$$

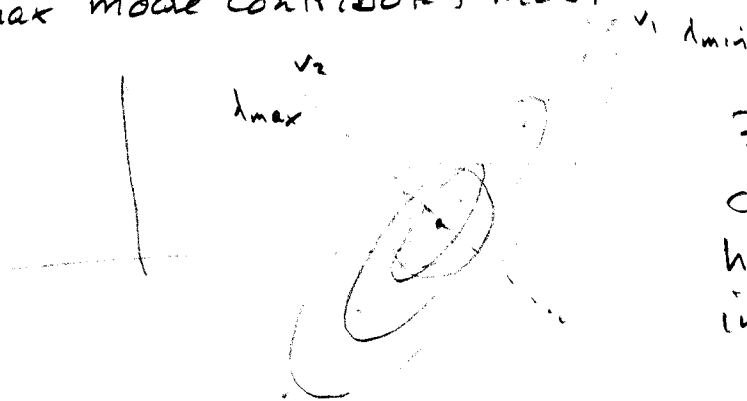
- Finally, we obtain the excess mse due to tap wander. We could do it like this:

$$J_{ex} = E[\epsilon^T R \epsilon] = \text{tr}[R E[\epsilon \epsilon^T]] \\ = \mu \frac{J_{min}}{2} \text{tr}[R] = \mu \frac{J_{min}}{2} \sum_{i=1}^M \lambda_i = \frac{1}{2} \mu J_{min} M r_u(0)$$

but it's more instructive to continue in the eigenvector basis

$$J_{ex} = E[\epsilon^T R \epsilon] = E[v^T Q^T R Q v] = E[v^T \Lambda v] \\ = \sum_{i=1}^M \lambda_i \sigma_{v_i}^2 \quad \text{so modes are weighted by eigenvals} \\ = \mu \frac{J_{min}}{2} \sum_{i=1}^M \lambda_i$$

The λ_{max} mode contributes most to J_{ex} . Why?



Because the circular blur climbs higher up the walls in the λ_{max} direction

In any case,

$$J_{ex} = \frac{1}{2} \mu J_{min} M r_u(0)$$

A surprisingly simple result! Next, we'll extract some meaning from it.

- Effect of step size parameter

- μ too small? slow convergence $n_c \approx \frac{1}{\mu \lambda_{\min}}$

- μ too big? $J_{\text{ex}} = \frac{\mu}{2} J_{\min} M r_u(0)$ too big

- An easy guideline: allow 10% excess mse

$$J_{\text{ex}} = 0.1 J_{\min}$$

$$\mu = \frac{0.2}{M r_u(0)} \quad \text{easy to set, since we can measure } r_u(0)$$

It doesn't sacrifice too much speed:

$$n_c \approx \frac{1}{\mu \lambda_{\min}} = \frac{M r_u(0)}{0.2 \lambda_{\min}} = \frac{\sum_i \lambda_i}{0.2 \lambda_{\min}}$$

$$\text{so} \quad \frac{\lambda_{\max}}{0.2 \lambda_{\min}} \leq n_c \leq \frac{M \lambda_{\max}}{0.2 \lambda_{\min}}$$

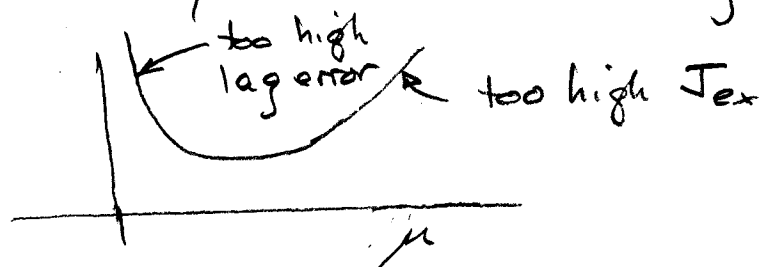
It still depends on the eigenvalue ratio.

- But we don't have to keep μ constant, which leads to another strategy:

gearshifting - start with large μ for initial convergence

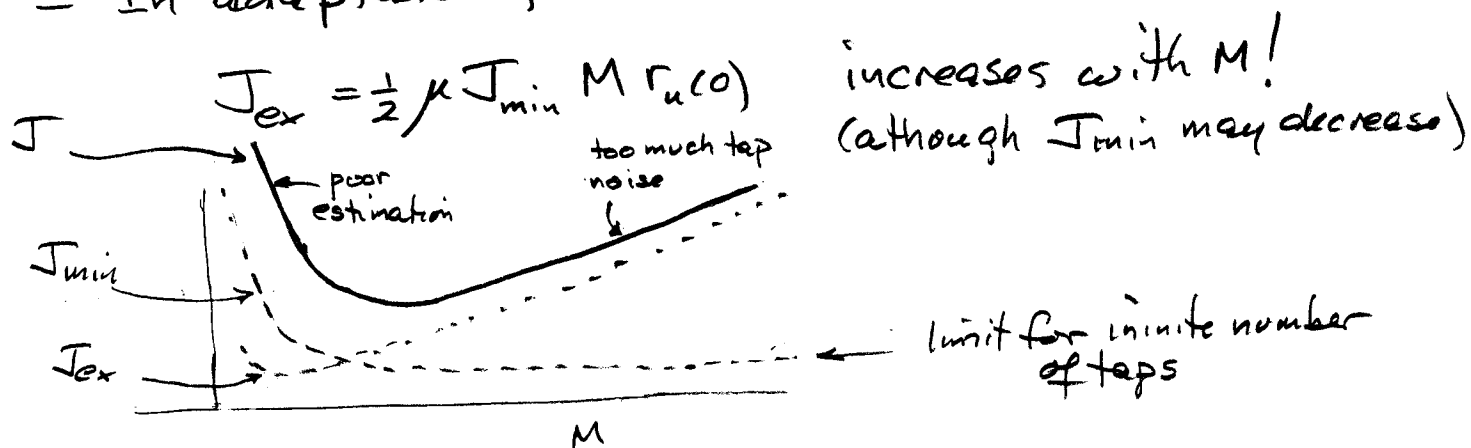
- then reduce it for small J_{ex}

- Tracking - if statistics drift, too small μ may cause filter to lag too much



• Effect of filter length

- Normally, we expect to improve performance if we increase filter length M — with diminishing returns, of course
- In adaptation, it's different:



- The growth with M is caused by the fact that every tap has same jitter variance $\mu \frac{J_{min}}{2}$ (same as σ_v^2 , circular) so additional taps add more misadjustment error.

Read Haykin 9.6, 9.7

Browse Haykin 9.11, especially (9.121), (9.122), (9.142) (9.144)