

6.5 Beating the Eigenvalue Spread - Decorrelation

6.5.1

- We have seen that the eigenvalue spread $\lambda_{\max}/\lambda_{\min}$ limits the performance. The slow mode ($\lambda_{\min}$) is r times slower than the fast mode ($\lambda_{\max}$), but if we try to speed it up by increasing μ , then the fast mode increases J_{ex} .
- Suggestion — if we could work individually with the modes, we could give them individual speed controls μ_i (step size parameter).

$$v_i(n+1) = (1 - \mu_i \lambda_i) v_i(n) + \mu_i z_i(n)$$

Specifically $\mu_i = \alpha / \lambda_i$

Then

$$v_i(n+1) = (1 - \alpha) v_i(n) + \alpha \frac{z_i(n)}{\lambda_i}$$

— Now modes are altered so all have same decay constant

$$n_c = \frac{-1}{\ln(1 - \alpha)} \approx \frac{1}{\alpha}$$

just as we did in steepest descent

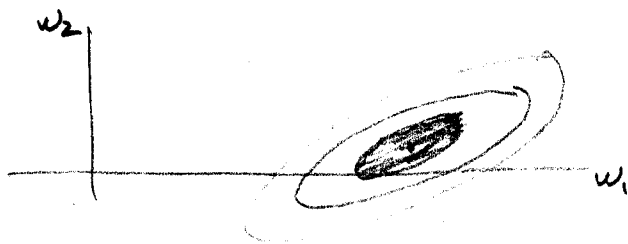
— Also, the perturbing noise variances become

$$E \left[\frac{|z_i(n)|^2}{\lambda_i^2} \right] = \frac{J_{\min} \lambda_i}{\lambda_i^2} = \frac{J_{\min}}{\lambda_i} \quad (\text{p. 6.4.4})$$

and the asymptotic jitter variance on mode i :

$$\sigma_{v_i}^2 = \frac{\alpha^2 J_{\min} / \lambda_i}{1 - (1 - \alpha)^2} \approx \frac{\alpha}{2} \frac{J_{\min}}{\lambda_i} \quad (\text{p. 6.4.5})$$

This means that the blur has same shape as contours



— So the effect on excess mse is

$$J_{\text{ex}} = \sum_{i=1}^M \lambda_i \sigma_{v_i}^2 = \frac{\alpha}{2} M J_{\text{min}}$$

- We have gained independence from eigenvals spread!

$$J_{\text{ex}} = \frac{\alpha}{2} M J_{\text{min}} \approx \frac{M J_{\text{min}}}{2 n_c}$$

For example, if we design for $J_{\text{ex}} = 0.1 J_{\text{min}}$

$$n_c = \frac{M}{0.2} = 5M$$

Compare this with original algorithm (p. 6.4.7)

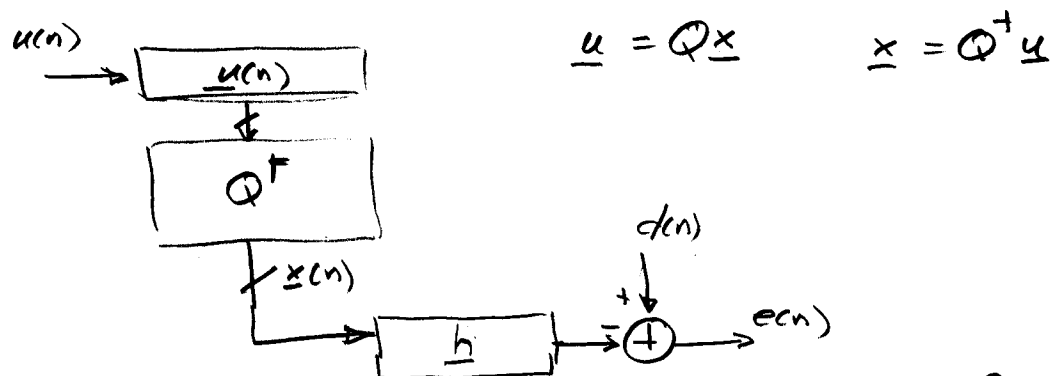
$$5r \leq n_c \leq 5rM \quad \text{where } r = \lambda_{\text{max}} / \lambda_{\text{min}}$$

- Where's the catch?

- We don't know R , so we don't know Q , so we can't decorrelate.

- But there are possibilities for adaptive decorrelation so we'll continue with decorrelated adaptation, assuming we know R

- Adaptation with explicit decorrelation in eigenspace: 6.5.3



Here we have performed transformation first, shift next, unlike previous derivation.

Correlation arrays: $R_x = E[\underline{x} \underline{x}^+] = Q^T R Q = \Lambda$

$$p_x = E[e^* \underline{x}] = Q^T p$$

- From p 6.4.2, update is

$$\begin{aligned} \underline{h}(n+1) &= \underline{h}(n) + \mu \underbrace{e^*(n) \underline{x}(n)}_{p_x - R_x \underline{h}(n) + \delta_x(n)} \\ &= (\mathbf{I} - \mu R_x) \underline{h}(n) + \mu p_x + \mu \delta_x(n) \end{aligned}$$

where $\delta_x(n) = e^*(n) \underline{x}(n) = Q^T e^*(n) \underline{u}(n) = \underline{v}(n)$ from earlier analysis p 6.4.4

$$\begin{aligned} \text{check: } E[\delta_x(n) \delta_x^+(n)] &= E[|e(n)|^2 \underline{x}(n) \underline{x}^+(n)] \approx J_{\min} R_x \\ &= J_{\min} \Lambda = E[\underline{v} \underline{v}^+] \text{ as expected} \end{aligned}$$

- Now do the shift. From the iteration above

$$\underline{h}_0 = R_x^{-1} p_x = Q^T R^{-1} Q Q^T p = Q^T R^{-1} p = Q^T \underline{w}_0$$

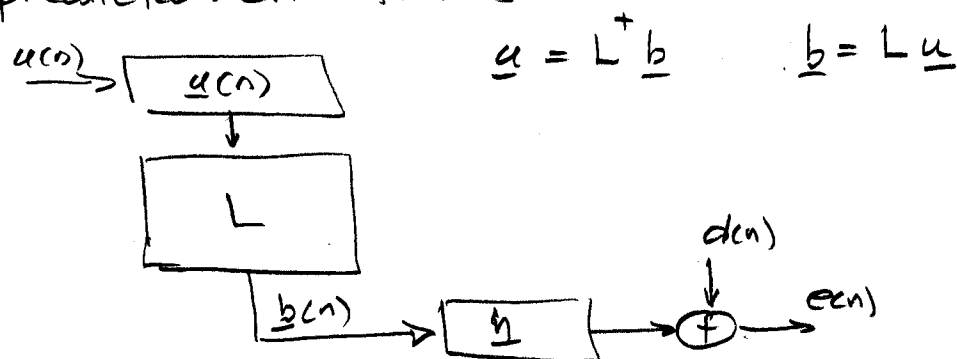
Then $\underline{h}(n) - \underline{h}_0 = \underline{v}(n)$ defined as in previous section

$$\text{and } \underline{v}(n+1) = (\mathbf{I} - \mu \Lambda) \underline{v}(n) + \mu \underline{v}(n)$$

- The matrix transformation gives wanted decoupling

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

- We can also decorrelate by Gram Schmidt, using prediction error filters



- Correlation arrays

$$R_b = E[\underline{b} \underline{b}^T] = L R L^T = D = \text{diag}[P_0, P_1, \dots, P_{M-1}]$$

$$P_b = E[e^* \underline{b}] = L P$$

- Iteration:

$$\underline{h}(n+1) = (I - \mu R_b) \underline{h}(n) + \mu P_b + \mu \underline{\gamma}_b(n)$$

where $\underline{\gamma}_b(n) = e^*(n) \underline{b}(n) = L e^*(n) u(n) = L \underline{\gamma}(n)$

and $E[\underline{\gamma}_b \underline{\gamma}_b^T] = L E[\underline{\gamma} \underline{\gamma}^T] L^T = J_{\min} L R L^T = J_{\min} D$

- New shift:

$$\underline{h}_0 = R_b^{-1} P_b = L R^{-1} L^T L P = L R^{-1} P = L \underline{w}_0$$

Define $\tilde{\underline{v}}(n) = \underline{h}(n) - \underline{h}_0$

and the iteration is

$$\begin{aligned} \tilde{\underline{v}}(n+1) &= (I - \mu R_b) \tilde{\underline{v}}(n) + \mu \underline{\gamma}_b(n) \\ &= (I - \mu D) \tilde{\underline{v}}(n) + \mu \underline{\gamma}_b(n) \end{aligned}$$

- Decouples to $\tilde{v}_i(n+1) = (1 - \mu P_i) \tilde{v}_i(n) + \mu \underbrace{\gamma_{b,i}(n)}_{\text{variance } P_i}$

so $\sigma_{\tilde{v}_i}^2 = \frac{\mu^2 J_{\min} P_i}{1 - (1 - \mu P_i)^2} \approx \mu \frac{J_{\min}}{2}$

$$J_{\text{ex}} = \sum_{i=0}^{M-1} P_i \sigma_{\tilde{v}_i}^2 = \mu \frac{J_{\min}}{2} \sum_{i=0}^{M-1} P_i$$

and since decoupled, we can choose

6.5.5

$$\mu_i = \frac{\alpha}{P_i}$$

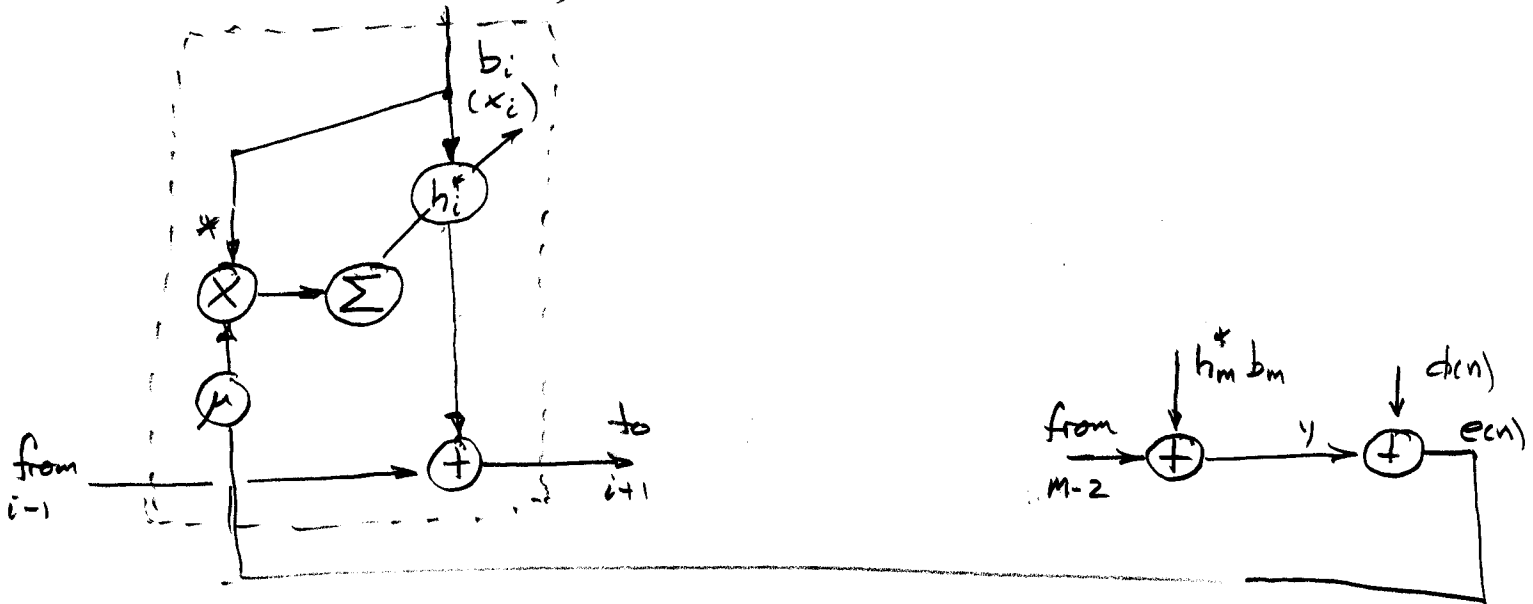
$$\tilde{v}_i(n+1) = (1-\alpha)\tilde{v}_i(n) + \alpha \frac{\tilde{d}_i(n)}{P_i} \quad n_c \approx \frac{1}{\alpha}$$

$$\sigma_{\tilde{v}_i}^2 \approx \frac{\alpha J_{min}}{2P_i}$$

$$J_{ex} = \frac{\alpha J_{min} M}{2}$$

- same advantages as in eigenspace decoupling:
 - independence from eigenvalue ratio
 - uniform convergence speed.

- In both cases, a stage looks like this



- Too bad we don't know R . But what if the lattice filter could be made adaptive? Then adaptive decorrelation.
 - Sets the stage for gradient adaptive lattice