

2.5 Linear MMSE Estimation

2.5.1

- Not all estimators are linear.

Not all mmse estimators are linear.

Requiring the estimator to be linear may cost accuracy.

But... combining linear estimation with mmse criterion has interesting properties:

- the linear coefficients and resulting variances depend only on second order statistics
- for Gaussian statistics, linear is optimum
 - mmse
 - MAP

- For Gaussian, we already have the estimator

$$\underline{\mu} = B D^{-1} \underline{x}$$

What's left to discover?

We'll derive linear mmse estimators for any statistics, Gaussian or not. More in Section

3 of the course, where we go through Haykin Chapt 5

2.5.1 Start with scalars

$$\hat{\theta} = h^* x, \quad e = \theta - \hat{\theta} \quad \text{assume zero means}$$

$$\text{minimize } \sigma_e^2 = E[(\theta - h^* x)(\theta^* - h x^*)]$$

differentiate w.r.t h

$$\frac{d\sigma_e^2}{dh}$$

not analytic ??

• Generalized derivative wrt complex variable

2.5.2

$$\frac{d}{dz} = \frac{1}{2} \left(\frac{\partial}{\partial u} - j \frac{\partial}{\partial v} \right) \quad z = u + jv$$

Try some obvious tests

$$\frac{dz}{dz} = \frac{1}{2} \left(\frac{\partial z}{\partial u} - j \frac{\partial z}{\partial v} \right) = \frac{1}{2} (1 - j \cdot j) = 1$$

$$\begin{aligned} \frac{dz^2}{dz} &= \frac{1}{2} \left(\frac{\partial z^2}{\partial u} - j \frac{\partial z^2}{\partial v} \right) = \frac{1}{2} \left(2z \frac{\partial z}{\partial u} - j 2z \frac{\partial z}{\partial v} \right) \\ &= 2z \end{aligned}$$

$$\frac{dz^*}{dz} = \frac{1}{2} \left(\frac{\partial z^*}{\partial u} - j \frac{\partial z^*}{\partial v} \right) = \frac{1}{2} (1 - j(-j)) = 0$$

Consistent with conventional derivative for any analytic function (C-R conditions $\frac{\partial f_r}{\partial u} = \frac{\partial f_i}{\partial v}$
 $\frac{\partial f_r}{\partial v} = -\frac{\partial f_i}{\partial u}$)

• Back to estimation:

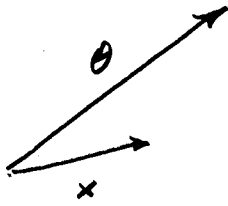
$$\begin{aligned} \frac{d\sigma_{\theta}^2}{dh} &= -E[(\theta - h^* x) x^*] = -\sigma_{\theta x}^2 + h^* \sigma_x^2 \\ &= -E[ex^*] = -\sigma_{ex}^2 \end{aligned}$$

This is the derivative, or gradient. Set it to zero.

$$h^* = \frac{\sigma_{\theta x}^2}{\sigma_x^2}$$

$$\sigma_{ex}^2 = 0 \quad \text{or the optimality condition}$$

$$\hat{\theta} = \frac{\sigma_{\theta x}^2}{\sigma_x^2} x$$



2.5.2 Multidimensional x

$$\hat{\theta} = \underline{h}^+ x \quad e = \theta - \underline{h}^+ x$$

$$\sigma_e^2 = E[(\theta - \underline{h}^+ x)(\theta^* - x^+ \underline{h})] = \sigma_\theta^2 - \underline{h}^+ \underline{p} - \underline{p}^+ \underline{h} + \underline{h}^+ R_x \underline{h}$$

$$\text{where } \underline{p} = E[\theta x^*], \quad R_x = E[x x^+]$$

Several components, so take derivative w.r.t $\underline{h}$

$$\frac{d\sigma_e^2}{d\underline{h}} = \begin{bmatrix} \frac{\partial \sigma_e^2}{\partial h_1} \\ \vdots \\ \frac{\partial \sigma_e^2}{\partial h_m} \end{bmatrix}$$

Read conventions in Haykin Appendix B carefully

$$= -E[(\theta - \underline{h}^+ x) x^*] = -\underline{p}^* + R_x^* \underline{h}^*$$

$$= -E[e x^*] = -\sigma_{ex}^2$$

conjugating, setting to zero

$$\left(\frac{d\sigma_e^2}{d\underline{h}}\right)^* = \underline{0} \Rightarrow R_x \underline{h} = \underline{p} \quad \sigma_{xe}^2 = \underline{0}$$

normal equations

orthogonality conditions

Equivalent conditions for optimum:

- The normal equations with covariance matrix of basis of estimate $\hat{\theta}$, and covariance vector of basis of estimate and quantity being estimated
- The estimation error is uncorrelated with (orthogonal to) basis of estimate

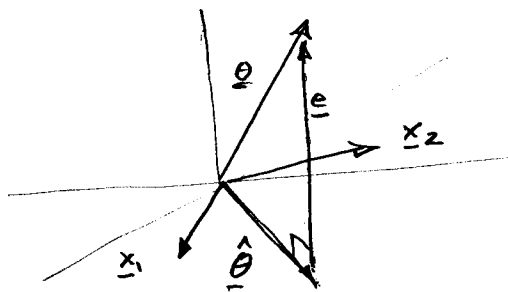
Minimized error:

$$\begin{aligned}
 \sigma_e^2 &= \sigma_\theta^2 - \underline{h}^T \underline{p} - \underline{p}^T \underline{h} + \underline{h}^T \underline{R}_x \underline{h} \quad \leftarrow \underline{h} = \underline{R}_x^{-1} \underline{p} \\
 &= \sigma_\theta^2 - \underline{p}^T \underline{R}_x^{-1} \underline{p} - \underline{p}^T \underline{R}_x^{-1} \underline{p} + \underline{p}^T \underline{R}_x^{-1} \underline{R}_x \underline{R}_x^{-1} \underline{p} \\
 &= \sigma_\theta^2 - \underline{p}^T \underline{R}_x^{-1} \underline{p} \\
 &= \sigma_\theta^2 - \sigma_\hat{\theta}^2 \quad \text{since } \sigma_\hat{\theta}^2 = E[|\hat{\theta}|^2] = E[\underline{h}^T \underline{x} \underline{x}^T \underline{h}] = \underline{p}^T \underline{R}_x^{-1} \underline{p}
 \end{aligned}$$

Relate to earlier $\underline{\mu} = \underline{B} \underline{D}^{-1} \underline{x}$ for Gaussian:

$$\begin{bmatrix} \theta \\ \underline{x} \end{bmatrix} \quad R = \begin{bmatrix} \sigma_\theta^2 & \underline{p}^T \\ \underline{p} & \underline{R}_x \end{bmatrix} \quad \hat{\theta} = \underline{h}^T \underline{x} = \underline{p}^T \underline{R}_x^{-1} \underline{x} = \underline{B} \underline{D}^{-1} \underline{x}$$

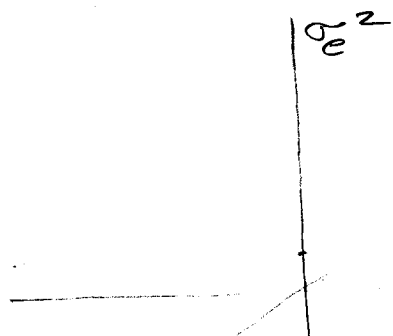
• Geometric view of multidimensional



- $\underline{x}_1, \underline{x}_2$ shown in plane
- and $\hat{\theta}$ lies in subspace spanned by them - a projection of $\hat{\theta}$ onto subspace
- the error $\underline{e}$ is orthogonal to the basis of estimate $\{\underline{x}_1, \underline{x}_2\}$, to $\hat{\theta}$ and to the entire subspace
- by Pythagoras $\sigma_e^2 = \sigma_\theta^2 - \sigma_{\hat{\theta}}^2$

The mse is a quadratic function of $\underline{h}$ (easier to sketch if real)

$$\sigma_e^2 = \sigma_\theta^2 - 2 \underline{p}^T \underline{h} + \underline{h}^T \underline{R}_x \underline{h}$$

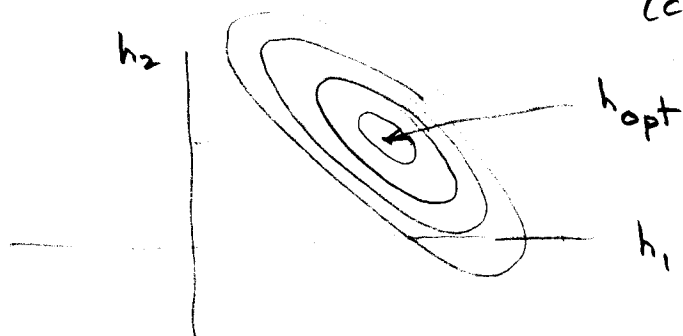


min σ_e^2
 $\underline{h}_{opt}$

h_1

we have the gradient of the error surface already

$$\nabla_{\underline{h}} \sigma_e^2 = -\underline{p} + R_x \underline{h} \quad \text{so can do steepest descent (conjugate gradient) iteration.}$$



top view, contour lines.
eigenvectors, eigenvalues?

• A convenient representation:

$$\underline{h} = \underline{h}_{opt} + \underline{\epsilon} = R_x^{-1} \underline{p} + \underline{\epsilon}$$

Then

$$\sigma_e^2 = \sigma_\theta^2 - (\underline{p}^T R_x^{-1} + \underline{\epsilon}^T) \underline{p} - \underline{p}^T (R_x^{-1} \underline{p} + \underline{\epsilon}) + (\underline{p}^T R_x^{-1} + \underline{\epsilon}^T) R_x (R_x^{-1} \underline{p} + \underline{\epsilon})$$

$$= \sigma_\theta^2 - \underline{p}^T R_x^{-1} \underline{p} + \underline{\epsilon}^T R_x \underline{\epsilon}$$

$$= \sigma_{e_{min}}^2 + \underline{\epsilon}^T R_x \underline{\epsilon} \quad \text{very simple!}$$

2.5.3 Multidimensional θ and $\underline{x}$

$$\hat{\theta} = H^T \underline{x} \quad \text{so} \quad H = [\underline{h}_1 | \underline{h}_2 | \dots | \underline{h}_K] \quad \underline{h}_i \text{ is column of coeffs to est } \theta_i$$

$$\underline{\epsilon} = \underline{\theta} - H^T \underline{x}$$

$$\text{minimize } J = E[\underline{\epsilon}^T \underline{\epsilon}] = \sum_{i=1}^K E[\epsilon_i^2] = \sum_{i=1}^K \sigma_{e_i}^2$$

The $\sigma_{e_i}^2$ are all non-negative, no opportunity for trade off among them, so decomposes to K separate optimizations

$$\min_{\underline{h}_i} \sigma_{e_i}^2 \Rightarrow \underline{h}_i = R_x^{-1} \underline{p}_i \quad \text{where } \underline{p}_i = E[\theta_i^* \underline{x}]$$

$$P = [\underline{p}_1 | \dots | \underline{p}_K] = E[\underline{x} \underline{\theta}^T]$$

Therefore $H = [h_1 | \dots | h_K] = R_x^{-1} [p_1 | p_2 | \dots | p_K]$
 $= R_x^{-1} P$

Relate to earlier $\mu = BD^{-1}x$ for Gaussian

$$\begin{bmatrix} \theta \\ x \end{bmatrix} \quad R = \left[\begin{array}{c|c} R_\theta & P^+ \\ \hline P & R_x \end{array} \right] \quad \hat{\theta} = H^+ x = P^+ R_x^{-1} x = BD^{-1}x$$

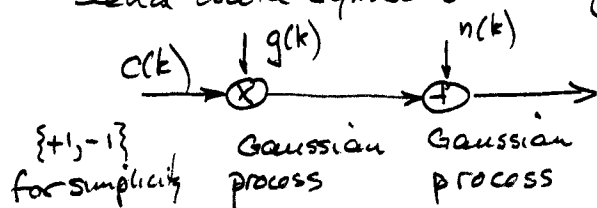
projection of entire subspace spanned by $\{\theta_i\}$ onto subspace spanned by $\{x_k\}$

Note mse can be written

$$\begin{aligned} J &= E[e^+ e] = E[(\theta^+ - x^+ H)(\theta - H^+ x)] = \sum_{i=1}^K \sigma_{e_i}^2 \\ &= E[\theta^+ \theta] - E[x^+ H \theta] - E[\theta^+ H^+ x] + E[x^+ H H^+ x] \\ &= \text{tr}[R_\theta] - \text{tr}[P^+ H] - \text{tr}[H^+ P] + \text{tr}[H^+ R_x H] \\ &= \text{tr}[R_\theta - P^+ H - H^+ P + H^+ R_x H] \end{aligned}$$

2.5.4 Example: Pilot Symbol Assisted Modulation

Send data symbols through a fading channel (e.g., mobile)



$$\text{SNR} = \sigma_g^2 / \sigma_n^2$$

Problem: $g(k)$ is complex Gaussian, so $c(k)$ could emerge with any phase, including reversal.

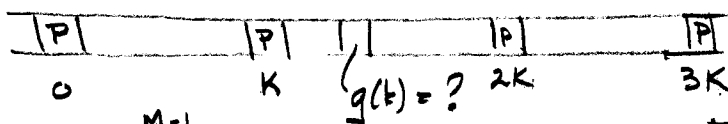
Solution: form an estimate $\hat{g}(k)$ of $g(k)$. Then do this:

$$r(k) \hat{g}^*(k) = \underbrace{g(k) \hat{g}^*(k)}_{\text{almost positive real}} c(k) + n(k) \hat{g}^*(k)$$

reversible

Here's how to get estimate $\hat{g}(k)$:

Make every K^{th} symbol a +1, $c(iK) = 1$ (pilot symbols)



$$\hat{g}(k) = \sum_{i=0}^{M-1} h_k^*(i) r(iK) \quad \text{or} \quad \hat{\theta}(k) = \underline{h}_k^T \underline{x}, \quad \theta = g(k)$$

where $\underline{x} = \begin{bmatrix} r(0) \\ r(K) \\ \vdots \\ r((M-1)K) \end{bmatrix}$

a linear interpolation
or linear estimate

the mmse estimate of $g(k)$ uses

$$\underline{h}_k = R_x^{-1} \underline{p}_k$$

$$R_x = \underline{x} \underline{x}^T$$

$$R_{x,ij} = R_g((i-j)K) + R_n((i-j)K)$$

autocorr. fns,
not matrices

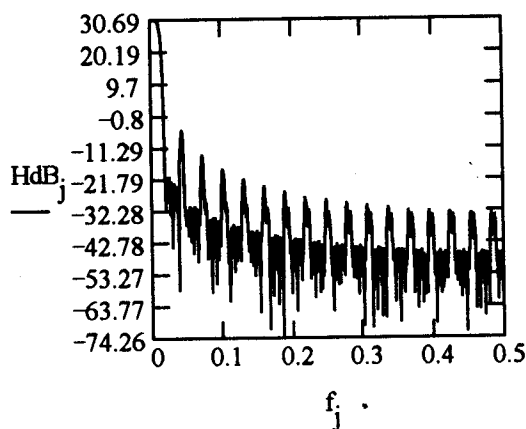
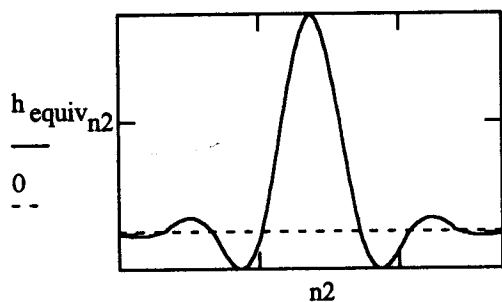
$$\underline{p}_k = E[g(k)^* \underline{x}] \quad [\underline{p}_k]_i = R_g(k-iK)$$

and the mmse estimate of all the $g(k)$:

$$\underline{H} = R_x^{-1} \underline{P}$$

$$\hat{\underline{\theta}} = \begin{bmatrix} \hat{g}(K+1) \\ \vdots \\ \hat{g}(2K+1) \end{bmatrix} = \underline{P}^T R_x^{-1} \underline{x}$$

Interesting fact — the coefficients can be considered as a polyphase filter; deinterleaved, it looks like this



$$\text{For } R_g(i) = \sigma_g^2 J_0(2\pi i f_0 T)$$

$$R_n(i) = \sigma_n^2 \delta(i)$$

Frequency response shows
periodic nulls characteristic
of interpolation filters

2.5.5 Partial Correlation

- Last topic on linear mmse estimation for a while — promise!
- Frequently, two or more quantities to be estimated are themselves correlated; e.g. $\begin{bmatrix} \theta_1 \\ \theta_2 \end{bmatrix}$ has R_θ

If they are both being estimated by $\underline{x}$, an interesting question is how much of their correlation was due to an $\underline{x}$ component — i.e., how correlated are they after removal of their projections on $\underline{x}$ subspace?

- Use tuples $\underline{\theta}_1$ has K_1 components, $\underline{\theta}_2$ has K_2 components, and $\underline{x}$ has M components.

$$\underline{e}_1 = \underline{\theta}_1 - \hat{\underline{\theta}}_1 = \underline{\theta}_1 - P^+ R_x^{-1} \underline{x} = \underline{\theta}_1 - E[\underline{\theta}_1 \underline{x}^+] R_x^{-1} \underline{x}$$

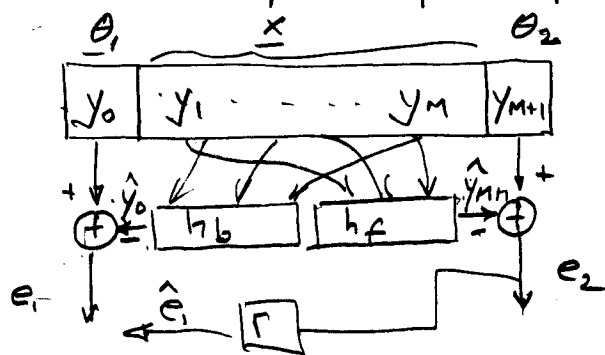
$$\underline{e}_2 = \underline{\theta}_2 - E[\underline{\theta}_2 \underline{x}^+] R_x^{-1} \underline{x}$$

Look at covariance matrix $E[\underline{e}_1 \underline{e}_2^+]$ $K_1 \times K_2$

Perhaps even project $\underline{e}_1$ onto $\underline{e}_2$

$$\hat{\underline{e}}_1 = \Gamma \underline{e}_2 \quad \Gamma = \underbrace{E[\underline{e}_1 \underline{e}_2^+] E[\underline{e}_2 \underline{e}_2^+]^{-1}}_{\text{PARCOR coefficients}}$$

- The PARCOR coefficients appear in forward and backward linear prediction (especially in speech work)



helps in forming
a recursive estimate

- If vectors are Gaussian, then $\underline{e}_1, \underline{e}_2$ related to conditional pdfs of y_0, y_{M+1} :

$$p_{\underline{e}_1}(\underline{e}_1) = p_{\underline{\theta}_1 | \underline{x}}(\underline{e}_1 + \hat{y}_0 | \underline{x}) \quad p_{\underline{e}_2}(\underline{e}_2) = p_{\underline{\theta}_2 | \underline{x}}(\underline{e}_2 + \hat{y}_{M+1} | \underline{x})$$

may be conditionally independent.

2.5.6 General View of Linear Estimation

- Really the last topic on linear estimation, quadratic criterion (I was just kidding last time).
- In an inner product space, we want to approximate $\underline{\theta}$ as a linear combination of basis vectors $\underline{x}_1, \dots, \underline{x}_M$ — that is, $\hat{\underline{\theta}}$ lies in subspace spanned by $\{\underline{x}_i\}$.

$$\hat{\underline{\theta}} = \sum_{i=1}^M h_i^* \underline{x}_i \quad \underline{e} = \underline{\theta} - \hat{\underline{\theta}}$$

Criterion — minimize length (norm) of $\underline{e}$
equivalently, squared norm.

Essential features: - error depends linearly on h coefficients
- criterion is quadratic function of error

- For any $\underline{\theta}$ in the subspace, we get coeffs of $\underline{\theta}$ by inner prod with successive $\underline{x}_k$:

$$(\underline{x}_k, \underline{\theta}) = \sum_{i=1}^M h_i (\underline{x}_k, \underline{x}_i) \quad k=1..M$$

Note: h_i lost its conjugate — see distributive rules for inner prod.

$$G \underline{h} = \underline{p} \quad \text{elementary linear algebra}$$

$$[G]_{k,i} = (\underline{x}_k, \underline{x}_i) \quad \text{Gram matrix of inner prods of basis functions with each other}$$

$$[\underline{p}]_k = (\underline{x}_k, \underline{\theta}) \quad \text{tuple of inner prods of basis functions with vector } \underline{\theta}$$

- But $\underline{\theta}$ may or may not be in the subspace, so we minimize

$$\|\underline{e}\|^2 = (\underline{e}, \underline{e}) = (\underline{\theta} - \hat{\underline{\theta}}, \underline{e}) = (\underline{\theta} - \hat{\underline{\theta}}, \underline{\theta} - \hat{\underline{\theta}}) \quad \text{①} \quad \text{②}$$

$$\text{Consider ①} \quad \|\underline{e}\|^2 = (\underline{\theta}, \underline{e}) - \sum_{i=1}^M h_i (\underline{x}_i, \underline{e})$$

$$\frac{d}{dh} \|\underline{e}\|^2 = - \begin{bmatrix} (\underline{x}_1, \underline{e}) \\ \vdots \\ (\underline{x}_M, \underline{e}) \end{bmatrix} = \underline{0}$$

Condition for min norm:
error is orthog to all basis vectors, hence to all members of the subspace.

Consider $\underline{e} = (\underline{\theta} - \sum_i h_i^* \underline{x}_i, \underline{\theta} - \sum_i h_i^* \underline{x}_i)$

$$\|\underline{e}\|^2 = (\underline{\theta}, \underline{\theta}) - \sum_i h_i^* (\underline{x}_i, \underline{\theta}) - \sum_i h_i (\underline{\theta}, \underline{x}_i) + \sum_i \sum_k h_i^* h_k (\underline{x}_i, \underline{x}_k)$$

Just for a change, find opt by

$$\frac{d}{dh^*} = - \begin{bmatrix} (\underline{x}_1, \underline{\theta}) \\ \vdots \\ (\underline{x}_m, \underline{\theta}) \end{bmatrix} + \begin{bmatrix} \sum_k h_k (\underline{x}_1, \underline{x}_k) \\ \vdots \\ \sum_k h_k (\underline{x}_m, \underline{x}_k) \end{bmatrix} = 0$$

or $-p + Gh = 0 \quad Gh = p$

This has same form as equations for h when $\underline{\theta}$ is in the space spanned by $\{\underline{x}_i\}$

$$p = \begin{bmatrix} (\underline{x}_1, \underline{\theta}) \\ \vdots \\ (\underline{x}_m, \underline{\theta}) \end{bmatrix}$$

in the space
(just a special case)

$$p = \begin{bmatrix} (\underline{x}_1, \underline{\theta}) \\ \vdots \\ (\underline{x}_m, \underline{\theta}) \end{bmatrix}$$

may not be in the space.

Finally the minimized $\|\underline{e}\|$. For any h ,

$$\|\underline{e}\|^2 = (\underline{\theta}, \underline{\theta}) - h^+ p - p^+ h + h^+ G h$$

Subst $h = G^{-1} p$, and

$$\|\underline{e}\|^2 = \|\underline{\theta}\|^2 - p^+ G^{-1} p$$

but $\|\hat{\underline{\theta}}\|^2 = (\sum_i h_i^* \underline{x}_i, \sum_i h_i^* \underline{x}_i) = h^+ G h \rightarrow p^+ G^{-1} p$ at optimum

so $\min \|\underline{e}\|^2 = \|\underline{\theta}\|^2 - \|\hat{\underline{\theta}}\|^2$

Pythagoras

