

1

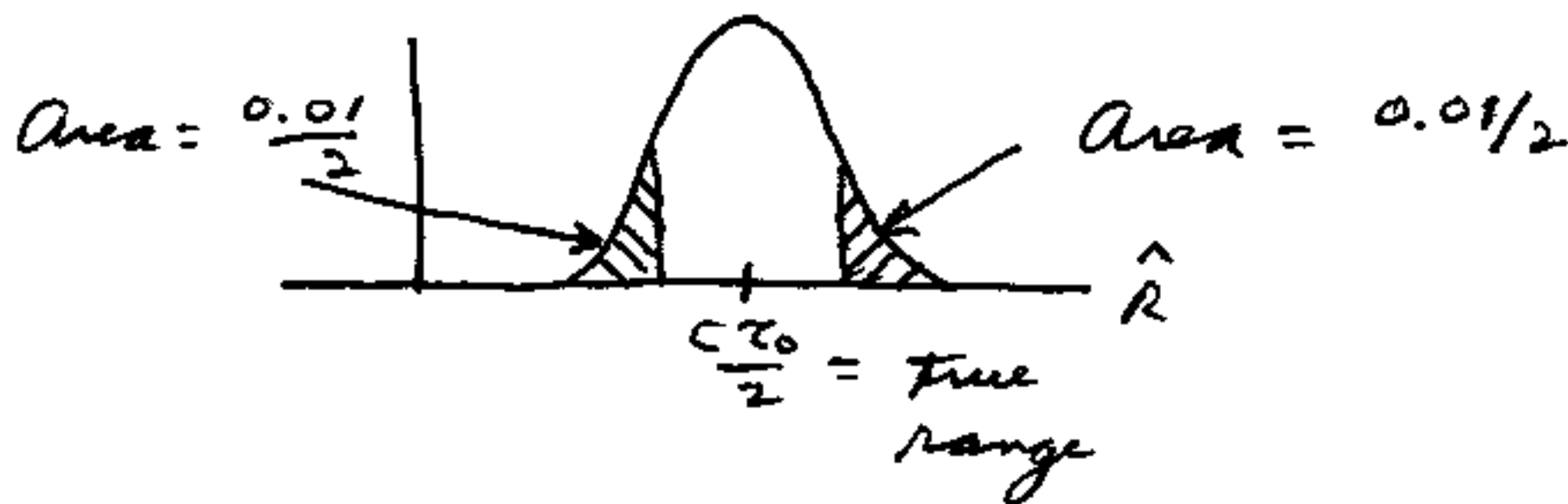
PROBLEM SOLUTIONS

FUNDAMENTALS OF STATISTICAL
SIGNAL PROCESSING: ESTIMATION
THEORY

BY STEVEN KAY

Chapter 1

- 1) Since $R = c\tau_0/2$, we use $\hat{R} = c\hat{\tau}_0/2$.
 The PDF is from $\hat{\tau}_0 \sim N(\tau_0, \sigma_{\hat{\tau}_0}^2)$,
 $\hat{R} \sim N(c\tau_0/2, \frac{c^2}{4} \sigma_{\hat{\tau}_0}^2)$



To be within 100 m we must have

$$Pr \{ |\hat{R} - c\tau_0/2| < 100 \} = 0.99$$

$$\Rightarrow Pr \left\{ \underbrace{\left| \frac{\hat{R} - c\tau_0/2}{\frac{c}{2} \sigma_{\hat{\tau}_0}} \right|}_{N(0,1)} < \frac{100}{\frac{c}{2} \sigma_{\hat{\tau}_0}} \right\} = 0.99$$

$$\Rightarrow \frac{100}{\frac{c}{2} \sigma_{\hat{\tau}_0}} = 2.58 \quad \text{or} \quad \sigma_{\hat{\tau}_0} = 2.6 \times 10^{-7} \text{ sec} \\ = 0.26 \mu\text{sec}$$

- 2) No, in fact θ could have been any value.

If θ were indeed 100, then

$$p(x; \theta) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x-100)^2}$$

and the probability of x being in the interval $[-97, 103]$ or $\mu \pm 3\sigma$ is 0.999.

Hence, his assertion is likely to be correct. However, we cannot be

certain, since if $\theta = 99$, then the probability of x being in the observed interval (for a single experiment) is

$$\int_{-97}^{103} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x-99)^2} dx =$$

$$\int_{-2}^4 \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}u^2} du = 0.977.$$

Thus, $\theta = 99$ is also highly likely.

3) $x = \theta + w$

$$p(x; \theta) = p_w(x - \theta)$$

For θ a random variable independent of w ,

$$p(x|\theta) = \frac{p_{x\theta}(x, \theta)}{p(\theta)} = \frac{p_{w\theta}(x - \theta, \theta)}{p(\theta)}$$

$$= \frac{p_w(x - \theta) p(\theta)}{p(\theta)} = p_w(x - \theta)$$

which is the same as before. If w and θ are not independent, then

$$p(x|\theta) = \frac{p_{w|\theta}(x - \theta) p(\theta)}{p(\theta)} = p_{w|\theta}(x - \theta)$$

which will be different than $p_w(x - \theta)$.

In general, $p(x; \theta) \neq p(x|\theta)$.

4) As shown in text $E(\hat{A}) = A$. Also,

$$E(\check{A}) = \frac{1}{N+2} (2A + (N-2)A + 2A) = A$$

Also, we know that $\text{var}(\hat{A}) = \sigma^2/N = 1/N$ and

$$\begin{aligned} \text{var}(\check{A}) &= \frac{1}{(N+2)^2} \left[4\sigma^2 + \sum_{n=1}^{N-2} \sigma^2 + 4\sigma^2 \right] \\ &= \frac{N+6}{(N+2)^2} \sigma^2 = \frac{N+6}{(N+2)^2} \end{aligned}$$

$$\begin{aligned} \text{var}(\check{A}) - \text{var}(\hat{A}) &= \frac{N+6}{(N+2)^2} - \frac{1}{N} \\ &= \frac{N(N+6) - (N+2)^2}{N(N+2)^2} \\ &= \frac{2N-4}{N(N+2)^2} > 0 \text{ for } N > 2 \end{aligned}$$

Hence, both estimators yield the correct value on the average but $\hat{A}$ has less variance. Conclusion is the same for any value of A .

5) $\hat{A}$ is not an estimator since to implement it requires knowledge of A (to determine the SNR).

Chapter 2

$$\begin{aligned}
 1) \quad E(\hat{\sigma}^2) &= E\left(\frac{1}{N} \sum_{n=0}^{N-1} x^2[n]\right) = \frac{1}{N} \sum_{n=0}^{N-1} E(x^2[n]) \\
 &= \sigma^2 \quad \text{for all } \sigma^2 > 0 \text{ (allowable values)} \\
 &\Rightarrow \text{unbiased}
 \end{aligned}$$

$$\begin{aligned}
 \text{var}(\hat{\sigma}^2) &= \frac{1}{N^2} \text{var}\left(\sum_n x^2[n]\right) \\
 &= \frac{1}{N^2} N \text{var}(x^2[n]) \quad (x[n]'s \text{ are IID} \\
 &\quad \Rightarrow x^2[n] \text{ are IID}) \\
 &= \frac{1}{N} \text{var}(x^2[n])
 \end{aligned}$$

$$\begin{aligned}
 \text{var}(x^2[n]) &= E(x^4[n]) - E(x^2[n])^2 \\
 &= 3\sigma^4 - \sigma^4 = 2\sigma^4
 \end{aligned}$$

$$\Rightarrow \text{var}(\hat{\sigma}^2) = 2\sigma^4/N \rightarrow 0 \text{ as } N \rightarrow \infty$$

Hence, the PDF of $\hat{\sigma}^2$ collapses about the true value as $N \rightarrow \infty$.

$$2) \quad \text{Let } \hat{\theta} = 2 \frac{1}{N} \sum_{n=0}^{N-1} x[n] \quad \text{since } E(x[n]) = \theta/2$$

$$E(\hat{\theta}) = \frac{2}{N} \sum_{n=0}^{N-1} \theta/2 = \theta$$

$$3) \quad \hat{A} = \frac{1}{N} \sum_{n=0}^{N-1} x[n] \quad \hat{A} \text{ is Gaussian since it is a } \underline{\text{linear}} \text{ function of independent Gaussian random variables. The mean was found to be } A \text{ and the variance is}$$

$$\begin{aligned}
 \text{var}(\hat{A}) &= \text{var}\left(\frac{1}{N} \sum_{n=0}^{N-1} x[n]\right) \\
 &= \frac{1}{N^2} \text{var}\left(\sum_{n=0}^{N-1} x[n]\right) \\
 &= \frac{1}{N^2} N \text{var}(x[n])
 \end{aligned}$$

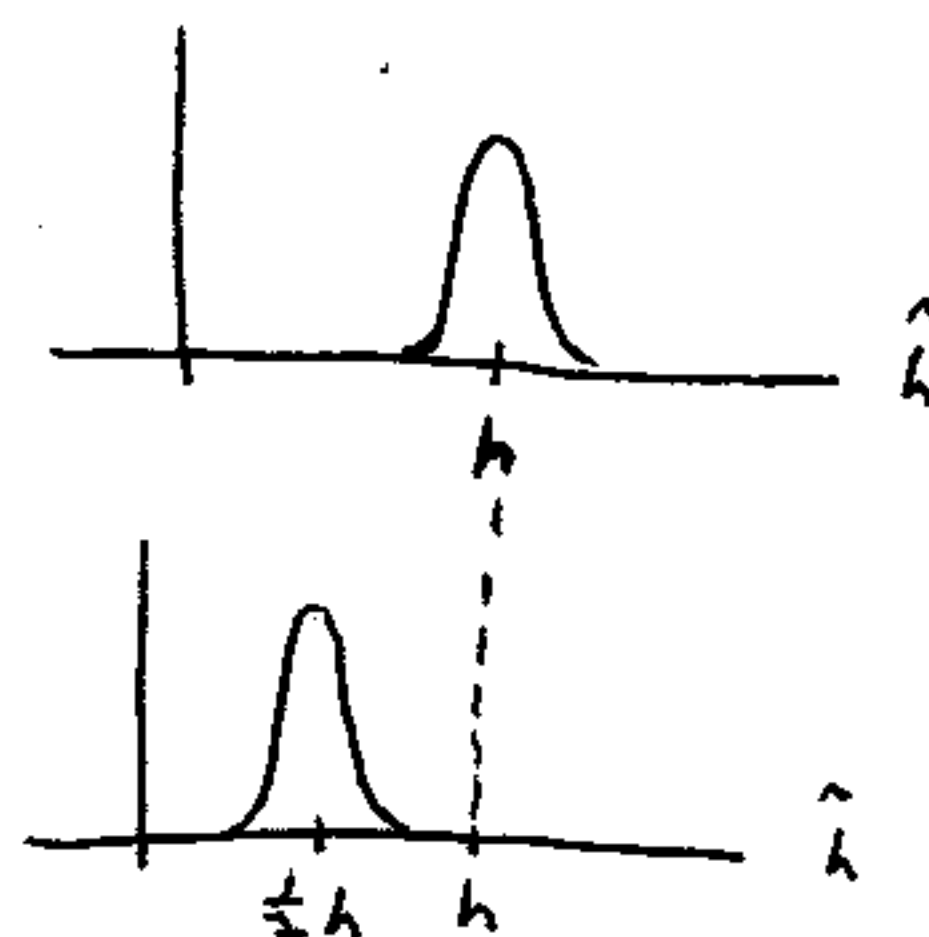
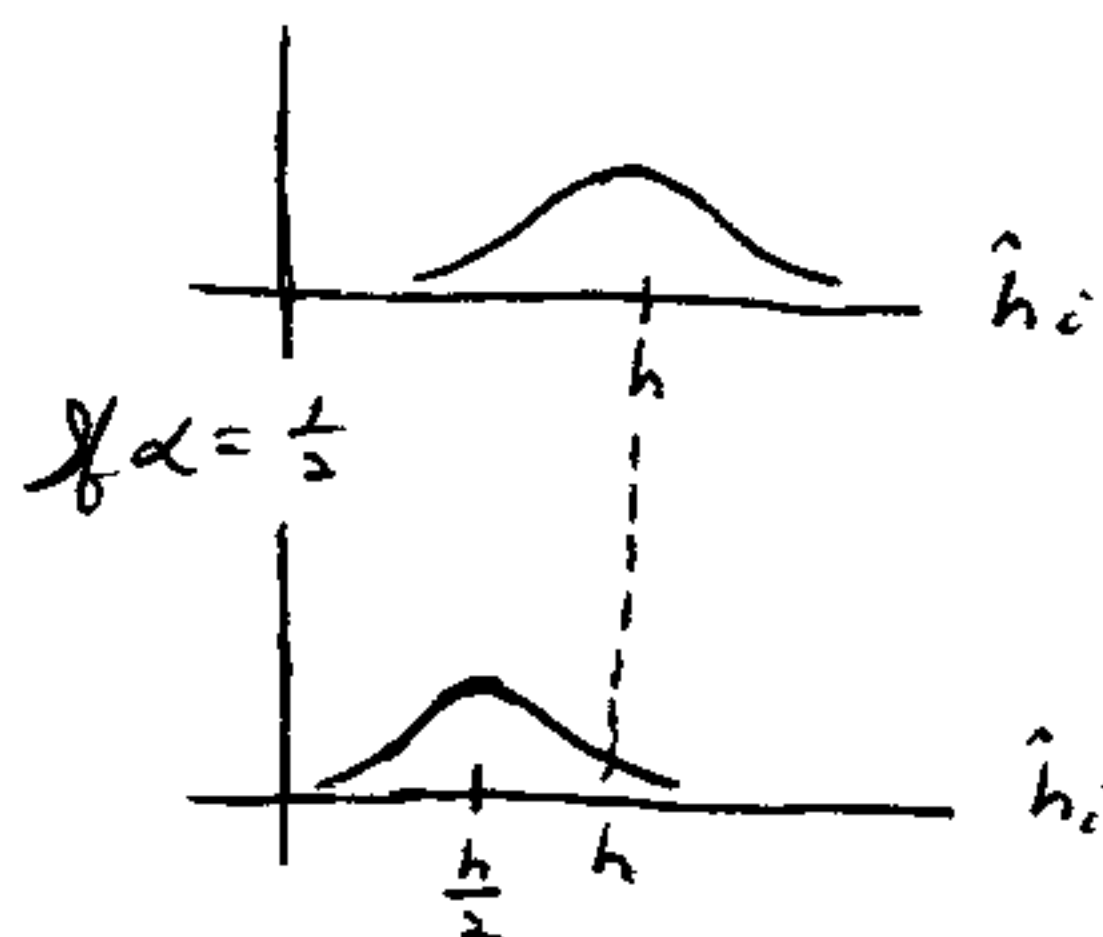
Since the $x[n]$'s are IID and thus uncorrelated

$$\Rightarrow \text{var}(\hat{A}) = \sigma^2/N.$$

$$4) \quad \hat{h} = \frac{1}{10} \sum_{i=1}^{10} h_i \quad E(\hat{h}) = \frac{1}{10} \sum_{i=1}^{10} E(h_i) = \alpha h$$

$$\text{var}(\hat{h}) = \text{var}(\hat{h}_i)/10 = 1/10$$

If $\alpha = 1$, we have



In second case ($\alpha = \frac{1}{2}$), averaging causes the PDF to be more heavily concentrated about the wrong value of h . The probability

of $\hat{h}$ being close to h actually decreases due to averaging. For $\alpha=1$, averaging, of course, is beneficial.

$$5) \quad X[N] \sim N(0, \sigma^2) \quad \text{or} \quad \frac{X[N]}{\sigma} \sim N(0, 1)$$

$$\Rightarrow \left(\frac{X[N]}{\sigma}\right)^2 \sim \chi^2_1 \quad \text{and}$$

$$y = \left(\frac{X[0]}{\sigma}\right)^2 + \left(\frac{X[1]}{\sigma}\right)^2 \sim \chi^2_2$$

$$\text{or } p(y) = \begin{cases} \frac{1}{2} e^{-y/2} & y > 0 \\ 0 & y < 0 \end{cases}$$

Transforming, we have $\hat{\sigma}^2 = \frac{\sigma^2}{2} y$
so that

$$p(\hat{\sigma}^2) = \frac{p_y(y(\hat{\sigma}^2))}{|d\hat{\sigma}^2/dy|}$$

$$= \frac{\frac{1}{2} e^{-\frac{1}{2}(2\hat{\sigma}^2/\sigma^2)}}{\sigma^2/2} \quad \begin{matrix} \hat{\sigma}^2 > 0 \\ 0 & \hat{\sigma}^2 < 0 \end{matrix}$$

$$= \frac{1}{\sigma^2} e^{-\hat{\sigma}^2/\sigma^2} \quad \begin{matrix} \hat{\sigma}^2 > 0 \\ 0 & \hat{\sigma}^2 < 0 \end{matrix}$$



Clearly, not symmetric.

$$\begin{aligned}
 \text{But } E(\hat{\sigma}^2) &= \int_0^\infty \hat{\sigma}^2 \frac{1}{\sigma^2} e^{-\hat{\sigma}^2/\sigma^2} d\hat{\sigma}^2 \\
 &= \sigma^2 \int_0^\infty u e^{-u} du \\
 &= \sigma^2 \left[-u e^{-u} - e^{-u} \right] \Big|_0^\infty \\
 &= \sigma^2
 \end{aligned}$$

$\hat{\sigma}^2$ is unbiased but PDF is not symmetric about σ^2 .

$$b) \quad E(\hat{A}) = \sum_{n=0}^{N-1} a_n A = A \Rightarrow \sum_{n=0}^{N-1} a_n = 1$$

$$\text{var}(\hat{A}) = \sum_{n=0}^{N-1} a_n^2 \text{var}(x[n]) = \sum_{n=0}^{N-1} a_n^2 \sigma^2$$

$$\text{Let } F = \sigma^2 \sum_{n=0}^{N-1} a_n^2 + \lambda \left(\sum_{n=0}^{N-1} a_n - 1 \right)$$

$$\frac{\partial F}{\partial a_i} = 2\sigma^2 a_i + \lambda = 0 \quad i = 0, 1, \dots, N-1$$

$$\Rightarrow a_i = -\lambda/2\sigma^2 \quad \text{for all } i$$

Thus, the a_i 's must be equal. But $\sum_{n=0}^{N-1} a_n = 1 \Rightarrow Na_i = 1$ or $a_i = 1/N$ or

$\hat{A}$ is just the sample mean estimator or as shown in Example 3.3, the MVU estimator.

$$7) \quad \frac{\hat{\theta} - \theta}{\sqrt{\text{var}(\hat{\theta})}} \sim N(0, 1) \quad \frac{\check{\theta} - \theta}{\sqrt{\text{var}(\check{\theta})}} \sim N(0, 1)$$

$$P_r \{ |\hat{\theta} - \theta| > \epsilon \} = P_r \left\{ \left| \frac{\hat{\theta} - \theta}{\sqrt{\text{var}(\hat{\theta})}} \right| > \frac{\epsilon}{\sqrt{\text{var}(\hat{\theta})}} \right\}$$

$$\text{Let } \Phi(x) = \int_{-\infty}^x \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}t^2} dt$$

= cumulative distribution
function for $N(0, 1)$

$$\Rightarrow P_r \{ |\hat{\theta} - \theta| > \epsilon \} = 2\Phi\left(\frac{-\epsilon}{\sqrt{\text{var}(\hat{\theta})}}\right)$$

$$\text{If } \text{var}(\hat{\theta}) < \text{var}(\check{\theta})$$

$$\Rightarrow \Phi\left(\frac{-\epsilon}{\sqrt{\text{var}(\hat{\theta})}}\right) < \Phi\left(\frac{-\epsilon}{\sqrt{\text{var}(\check{\theta})}}\right)$$

$$\text{or } P_r \{ |\hat{\theta} - \theta| > \epsilon \} < P_r \{ |\check{\theta} - \theta| > \epsilon \}$$

$$8) \text{ From Prob. 2.3 } \hat{A} \sim N(A, \sigma^2/N)$$

$$P_r \{ |\hat{A} - A| > \epsilon \} = P_r \left\{ \left| \frac{\hat{A} - A}{\sqrt{\sigma^2/N}} \right| > \frac{\epsilon}{\sqrt{\sigma^2/N}} \right\}$$

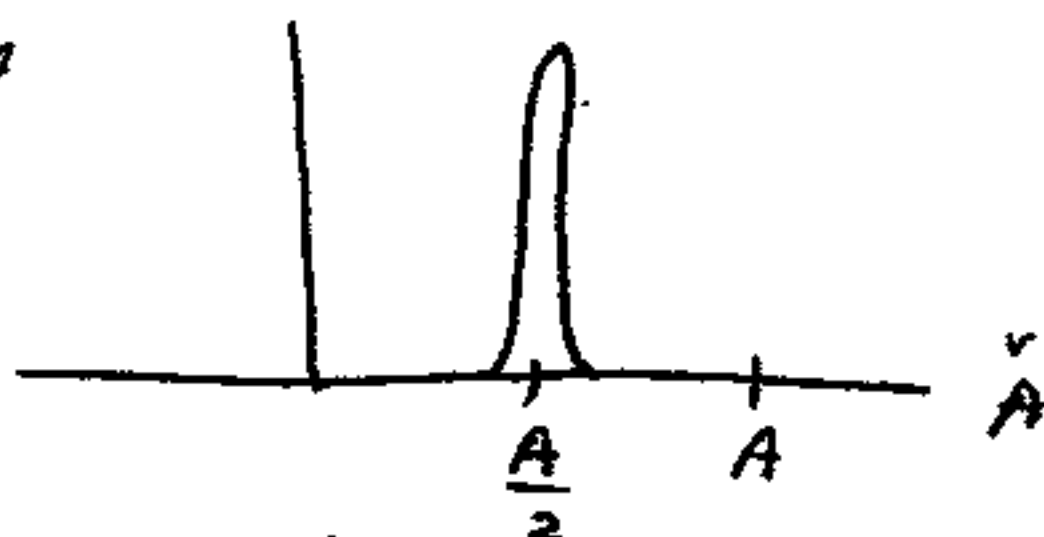
$$= 2\Phi\left(\frac{-\epsilon}{\sqrt{\sigma^2/N}}\right) \rightarrow 0$$

$$\text{Since } \frac{-\epsilon}{\sqrt{\sigma^2/N}} \rightarrow -\infty \text{ as } N \rightarrow \infty$$

$$\text{If } \hat{A} = \frac{1}{2N} \sum_{n=0}^{N-1} x[n] \text{ is used,}$$

$$\hat{A} \sim N(A/2, \sigma^2/4N)$$

As $N \rightarrow \infty$, the variance $\rightarrow 0$ and the PDF approaches



Thus, $\lim_{N \rightarrow \infty} P\{| \hat{A} - A | > \epsilon\} = 0$ for any $\epsilon > 0$

$\hat{A}$ is consistent while $\check{A}$ is inconsistent.

9) $\hat{\theta} = (\hat{A})^2$ where $\hat{A} \sim N(A, \sigma^2/N)$
 $E(\hat{\theta}) = E\{\hat{A}^2\} = \text{var}(\hat{A}) + E(\hat{A})^2$
 $= \sigma^2/N + A^2 = \theta + \sigma^2/N \neq \theta$

$\hat{\theta}$ is biased but as $N \rightarrow \infty$, it is unbiased. $\hat{\theta}$ is said to be asymptotically unbiased (for large data records).

10) Clearly, $E(\hat{A}) = A$. To find $E(\hat{\theta}^2)$

$$E(\hat{\theta}^2) = \frac{1}{N-1} \sum_{n=0}^{N-1} E\{(x[n] - \hat{A})^2\}$$

$$\text{But } E\{(x[n] - \hat{A})^2\} = E\left\{\left(x[n] - \frac{1}{N} \sum_{m=0}^{N-1} x[m]\right)^2\right\}$$

$$= E\left\{\left(x[n]\left(1 - \frac{1}{N}\right) - \frac{1}{N} \sum_{\substack{m=0 \\ m \neq n}}^{N-1} x[m]\right)^2\right\}$$

$$= \left(\frac{N-1}{N}\right)^2 E\{x^2[n]\} - 2 \frac{(N-1)}{N^2} E\left\{x[n] \sum_{\substack{m=0 \\ m \neq n}}^{N-1} x[m]\right\}$$

$$+ \frac{1}{N^2} E \left[\left(\sum_{\substack{m=0 \\ m \neq n}}^{N-1} x[m] \right)^2 \right]$$

$$= \left(\frac{N-1}{N} \right)^2 (\sigma^2 + A^2) - \frac{2(N-1)}{N^2} E[x[n]] E \left(\sum_{\substack{m=0 \\ m \neq n}}^{N-1} x[m] \right) \\ + \frac{1}{N^2} \left[\text{var} \left(\sum_{\substack{m=0 \\ m \neq n}}^{N-1} x[m] \right) + E \left(\sum_{\substack{m=0 \\ m \neq n}}^{N-1} x[m] \right)^2 \right]$$

$$= \left(\frac{N-1}{N} \right)^2 (\sigma^2 + A^2) - 2 \frac{(N-1)}{N^2} A (N-1) A \\ + \frac{1}{N^2} (N-1) \sigma^2 + \frac{1}{N^2} [(N-1) A]^2$$

$$= \sigma^2 \left[\frac{N^2 - 2N + 1 + N - 1}{N^2} \right] = \sigma^2 \frac{N-1}{N}$$

$$\Rightarrow E(\hat{\sigma}^2) = \frac{1}{N-1} \sum_{n=0}^{N-1} \sigma^2 \frac{N-1}{N} = \sigma^2$$

$\Rightarrow \hat{\sigma}^2$ is unbiased.

$$11) \quad \hat{\theta} = g(x[0])$$

$$E(\hat{\theta}) = \theta \Rightarrow \int g(x[0]) p(x[0]) dx[0] = \theta$$

$$\text{or } \int_0^{1/\theta} g(x[0]) \theta dx[0] = \theta$$

$$\int_0^{1/\theta} g(u) du = 1 \quad \text{for all } \theta > 0$$

Now suppose a g could be found. Then

for any $\theta_2 < \theta_1$, we would have

$$\int_0^{\theta_2} g(u) du = 1$$

$$\int_0^{\theta_1} g(u) du = 1$$

and subtracting the two gives

$$\int_{\theta_2}^{\theta_1} g(u) du = 0 \quad \text{for any } \theta_2 < \theta_1$$

Clearly, we must have $g(u) = 0$ for all u , which produces a biased estimator.

Chapter 3

$$1) \quad p(x(n); \theta) = \frac{1}{\theta} (u(x(n)) - u(x(n) - \theta))$$

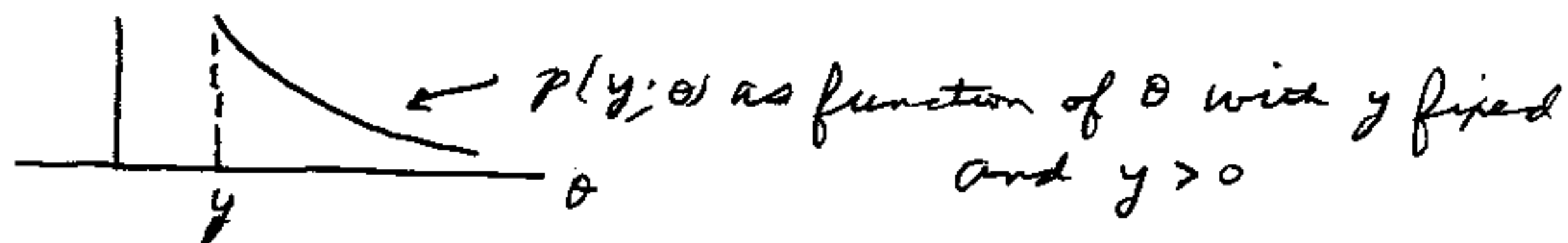
$$\text{where } u(x) = \begin{cases} 1 & x > 0 \\ 0 & x \leq 0 \end{cases}$$

Since $p(x; \theta) = \prod_{n=0}^{N-1} p(x(n); \theta)$, it is enough to show that

$$E \left[\frac{\partial \text{LN} p(x(n); \theta)}{\partial \theta} \right] \neq 0$$

(The expectation will be independent of n).
Let $y = x(n)$ so that

$$p(y; \theta) = \frac{1}{\theta} (u(y) - u(y - \theta))$$



For $\theta > y$

$$E \left[\frac{\partial \text{LN} p(y; \theta)}{\partial \theta} \right] = E \left[\frac{\partial \text{LN} 1/\theta}{\partial \theta} \right] \\ = -1/\theta \neq 0$$

$$2) \quad X(0) = A + W(0)$$

$$p_x(x(0); A) = p_w(x(0) - A) = p(x(0) - A)$$

$$I(A) = E \left[\left(\frac{\partial \text{LN} p(x(0) - A)}{\partial A} \right)^2 \right]$$

$$\begin{aligned}
&= E \left[\left(\frac{\partial \ln p(x|0)-A}{\partial (x|0)-A} (-1) \right)^2 \right] \\
&= E \left[\left(\frac{1}{p(x|0)-A} \frac{\partial p(x|0)-A}{\partial (x|0)-A} \right)^2 \right] \\
&= \int_{-\infty}^{\infty} \left(\frac{\partial p(x|0)-A}{\partial (x|0)-A} \right)^2 \frac{1}{p^2(x|0)-A} p_x(x|0; A) dx|0 \\
&= \int_{-\infty}^{\infty} \left(\frac{\partial p(x|0)-A}{\partial (x|0)-A} \right)^2 \frac{1}{p^2(x|0)-A} p(x|0)-A dx|0
\end{aligned}$$

Letting $u \equiv x|0)-A$

$$I(A) = \int_{-\infty}^{\infty} \frac{\left(\frac{dp(u)}{du} \right)^2}{p(u)} du$$

For $p(u) = \frac{1}{\sqrt{2}\sigma} e^{-\sqrt{2}|u|/\sigma}$ which is even in u

$$dp/du = -\frac{\sqrt{2}}{\sigma} \frac{1}{\sqrt{2}\sigma} e^{-\sqrt{2}u/\sigma} \quad u > 0$$

$$\begin{aligned}
I(A) &= 2 \int_0^{\infty} \frac{\frac{1}{\sigma^4} e^{-2\sqrt{2}u/\sigma}}{\frac{1}{\sqrt{2}\sigma} e^{-\sqrt{2}u/\sigma}} du \\
&= \frac{2\sqrt{2}\sigma}{\sigma^4} \int_0^{\infty} e^{-\sqrt{2}u/\sigma} du = \frac{2\sqrt{2}\sigma}{\sigma^4} \frac{\sqrt{2}\sigma}{2} \\
&= 2/\sigma^2
\end{aligned}$$

$\Rightarrow \text{var}(\hat{A}) \geq \sigma^2/2$ The CRLB is half of that for the Gaussian case. In fact,

it can be shown that the Gaussian PDF produces the largest CRLB.

$$3) \quad p(x; A) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_{n=0}^{N-1} (x(n) - Ar^n)^2}$$

$$\begin{aligned} \frac{\partial \text{LNP}}{\partial A} &= -\frac{2}{2\sigma^2} \sum_n (x(n) - Ar^n) (-r^n) \\ &= \frac{1}{\sigma^2} \sum_n (x(n) - Ar^n) r^n \end{aligned}$$

$$\frac{\partial^2 \text{LNP}}{\partial A^2} = -\frac{1}{\sigma^2} \sum_n r^{2n}$$

$$-E\left[\frac{\partial^2 \text{LNP}}{\partial A^2}\right] = \frac{1}{\sigma^2} \sum_n r^{2n}$$

$$\text{or } \text{var}(\hat{A}) \geq \frac{\sigma^2}{\sum_{n=0}^{N-1} r^{2n}} \quad (\text{or use (3.14)})$$

To show that an efficient estimator exists

$$\begin{aligned} \frac{\partial \text{LNP}}{\partial A} &= \frac{1}{\sigma^2} \left(\sum_n x(n) r^n - A \sum_n r^{2n} \right) \\ &= \underbrace{\frac{\sum_n r^{2n}}{\sigma^2}}_{I(A)} \left(\underbrace{\frac{\sum_n x(n) r^n}{\sum_n r^{2n}}}_{\hat{A}} - A \right) \end{aligned}$$

$\hat{A}$ is efficient and $1/I(A)$ is its variance.

$$\text{var}(\hat{A}) \rightarrow \sigma^2(1-r^2) \quad 0 < r < 1$$

$$\rightarrow 0$$

$$r \geq 1$$

as $N \rightarrow \infty$.

$$5) \quad p(\underline{x}; r) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_{n=0}^{N-1} (x[n] - r^n)^2}$$

$$\frac{\partial \text{LNP}}{\partial r} = \frac{1}{\sigma^2} \sum_n (x[n] - r^n) n r^{n-1}$$

Can't put into form $I(r)(\hat{r} - r) \Rightarrow$
no efficient estimator.

$$\frac{\partial^2 \text{LNP}}{\partial r^2} = \frac{1}{\sigma^2} \sum_n \left[x[n] n(n-1) r^{n-2} - n(2n-1) r^{2n-2} \right]$$

$$\begin{aligned} E \left[\frac{\partial^2 \text{LNP}}{\partial r^2} \right] &= \frac{1}{\sigma^2} \sum_n \left[n(n-1) r^{2n-2} - n(2n-1) r^{2n-2} \right] \\ &= -\frac{1}{\sigma^2} \sum_n n^2 r^{2n-2} \end{aligned}$$

$$\text{var}(\hat{r}) \geq \frac{\sigma^2}{\sum_{n=0}^{N-1} n^2 r^{2n-2}} \quad (\text{could also use (3.14)})$$

$$5) \quad p(\underline{x}; A) = \frac{1}{(2\pi)^{N/2} \det(\underline{C})^{1/2}} e^{-\frac{1}{2} (\underline{x} - A\underline{1})^T \underline{C}^{-1} (\underline{x} - A\underline{1})}$$

$$\text{where } \underline{1} = [1 \ 1 \ \dots \ 1]^T$$

$$\frac{\partial \text{LNP}}{\partial A} = -\frac{1}{2} \frac{\partial}{\partial A} \left[\underline{x}^T \underline{C}^{-1} \underline{x} - 2 \underline{1}^T \underline{C}^{-1} \underline{x} A + \underline{1}^T \underline{C}^{-1} \underline{1} A^2 \right]$$

$$\begin{aligned}
 &= \underline{1}^T \underline{C}^{-1} \underline{x} - \underline{1}^T \underline{C}^{-1} \underline{1} A \\
 &= \underbrace{\underline{1}^T \underline{C}^{-1} \underline{1}}_{I(A)} \left(\underbrace{\frac{\underline{1}^T \underline{C}^{-1} \underline{x}}{\underline{1}^T \underline{C}^{-1} \underline{1}}}_{\hat{A}} - A \right)
 \end{aligned}$$

$\hat{A}$ is efficient and has variance $1/I(A)$.

6)

$$\begin{aligned}
 X(0) &\sim N(\theta, 1) \\
 X(1) &\sim N(\theta, 1) \quad \theta \geq 0 \\
 &\quad N(\theta, 2) \quad \theta < 0
 \end{aligned}$$

$$\begin{aligned}
 p(\underline{x}; \theta) &= \frac{1}{2\pi} e^{-\frac{1}{2} [(x(0)-\theta)^2 + (x(1)-\theta)^2]} \quad \theta \geq 0 \\
 &\quad \frac{1}{2\pi\sqrt{2}} e^{-\frac{1}{2} [(x(0)-\theta)^2 + \frac{1}{2}(x(1)-\theta)^2]} \quad \theta < 0
 \end{aligned}$$

For $\theta \geq 0$

$$\begin{aligned}
 \frac{\partial \ln p}{\partial \theta} &= -\frac{1}{2} [2(x(0)-\theta)(-1) + 2(x(1)-\theta)(-1)] \\
 &= (x(0)-\theta) + (x(1)-\theta)
 \end{aligned}$$

$$\frac{\partial^2 \ln p}{\partial \theta^2} = -2 \Rightarrow E\left(-\frac{\partial^2 \ln p}{\partial \theta^2}\right) = 2$$

$$\text{var}(\hat{\theta}) \geq \frac{1}{2}$$

For $\theta < 0$

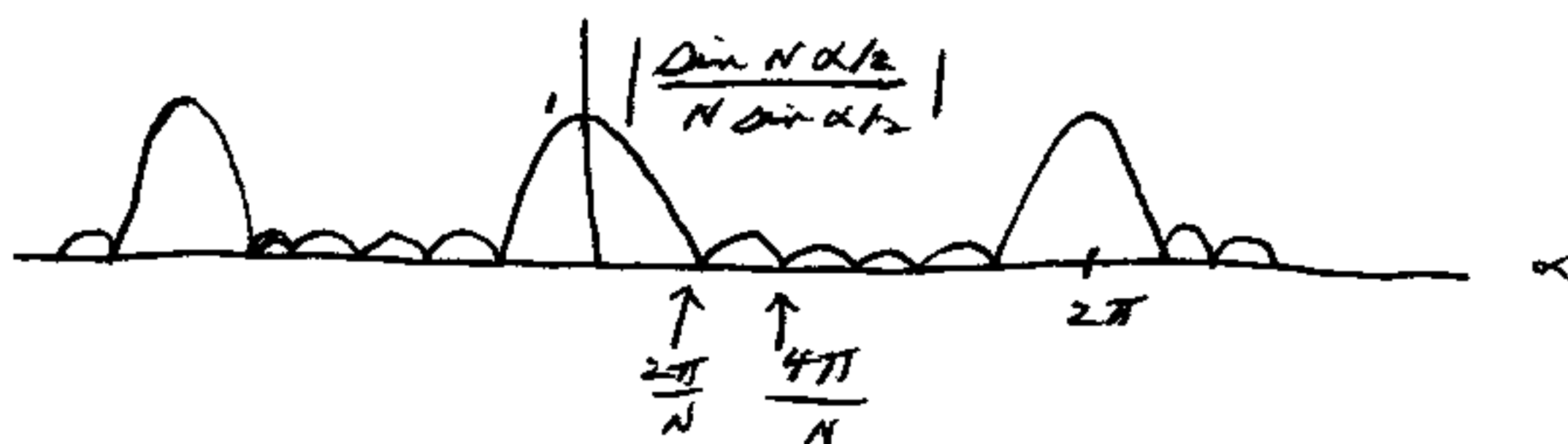
$$\begin{aligned}
 \frac{\partial \ln p}{\partial \theta} &= -\frac{1}{2} [-2(x(0)-\theta) - (x(1)-\theta)] \\
 &= (x(0)-\theta) + \frac{1}{2}(x(1)-\theta)
 \end{aligned}$$

$$\frac{\partial^2 LNP}{\partial \theta^2} = -3/2 \Rightarrow E\left[-\frac{\partial^2 LNP}{\partial \theta^2}\right] = 3/2$$

$$\text{var}(\hat{\theta}) \geq 2/3$$

7) Let $\alpha = 4\pi f_0$, $\beta = 2\phi$

$$\begin{aligned} \frac{1}{N} \sum_{n=0}^{N-1} \cos(\alpha n + \beta) &= \frac{1}{N} \text{Re} \left[\sum_n e^{j(\alpha n + \beta)} \right] \\ &= \frac{1}{N} \text{Re} \left[e^{j\beta} \frac{1 - e^{j\alpha N}}{1 - e^{j\alpha}} \right] \\ &= \frac{1}{N} \text{Re} \left[e^{j\beta} \frac{e^{j\alpha N/2}}{e^{j\alpha/2}} \frac{e^{-j\alpha N/2} - e^{j\alpha N/2}}{e^{-j\alpha/2} - e^{j\alpha/2}} \right] \\ &= \frac{1}{N} \text{Re} \left[e^{j\beta} e^{j\alpha(N-1)/2} \frac{\sin N\alpha/2}{\sin \alpha/2} \right] \\ &= \frac{\sin N\alpha/2}{N \sin \alpha/2} \cos \left[\alpha \left(\frac{N-1}{2} \right) + \beta \right] \end{aligned}$$



As long as α is not near 0 or 2π , this term is approximately zero. But $\alpha = 4\pi f_0 \Rightarrow f_0$ cannot be near 0 or $1/2$.

$$8) \quad p(\underline{x}; A) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_{n=0}^{N-1} (x[n] - A)^2}$$

$$\frac{\partial \ln p}{\partial A} = \frac{1}{\sigma^2} \sum_n (x[n] - A)$$

$$\left(\frac{\partial \ln p}{\partial A} \right)^2 = \frac{1}{\sigma^4} \sum_m \sum_n (x[m] - A)(x[n] - A)$$

$$\begin{aligned} E \left[\left(\frac{\partial \ln p}{\partial A} \right)^2 \right] &= \frac{1}{\sigma^4} \sum_m \sum_n \underbrace{E[(x[m] - A)(x[n] - A)]}_{\sigma^2 \delta_{mn}} \\ &= \frac{1}{\sigma^4} \sum_n \sigma^2 = N/\sigma^2 \end{aligned}$$

$$\text{var}(\hat{A}) \geq \sigma^2/N$$

9) Using (3.32)

$$\mathbf{I}(A) = \left(\frac{\partial \underline{\mu}(A)}{\partial A} \right)^T \mathbf{C}^{-1} \frac{\partial \underline{\mu}(A)}{\partial A}$$

$$\underline{\mu}(A) = [A \ A]^T \Rightarrow \frac{\partial \underline{\mu}(A)}{\partial A} = \underline{1}$$

$$\text{var}(\hat{A}) \geq \frac{1}{\underline{1}^T \mathbf{C}^{-1} \underline{1}} \quad (\text{or use approach of Prob 3.5})$$

$$\mathbf{C}^{-1} = \frac{1}{\sigma^2} \frac{\begin{bmatrix} 1 & -\rho \\ -\rho & 1 \end{bmatrix}}{1 - \rho^2}$$

$$\Rightarrow \text{var}(\hat{A}) \geq \frac{\sigma^2(1 - \rho^2)}{2 - 2\rho} = \frac{\sigma^2}{2} (1 + \rho)$$

If $\rho = 0$, $\text{var}(\hat{A}) \geq \sigma^2/2$, as expected.

If $\rho \rightarrow 1$, $\text{var}(\hat{A}) \geq \sigma^2$. This is the same bound as for one sample and occurs because as $\rho \rightarrow 1$, $W[0]$ and $W[1]$ will be equal. Hence, we have only one independent data sample. If $\rho \rightarrow -1$, $\text{var}(\hat{A}) \geq 0$ and in fact, in this case $W[0] = -W[1]$. Thus,

$$\begin{aligned}\hat{A} &= \frac{1}{2}(X[0] + X[1]) \\ &= \frac{1}{2}(A + W[0] + A - W[0]) = A\end{aligned}$$

for any realization of the noise samples. Additivity property of Fisher information only holds for independent samples. In this example we could have

$$i(A) \leq \mathcal{I}(A) < \infty i(A)$$

where $i(A) = 1/\sigma^2$.

$$10) \quad [\mathcal{I}(\underline{\theta})]_{ij} = E \left[\frac{\partial \text{LNP}}{\partial \theta_i} \frac{\partial \text{LNP}}{\partial \theta_j} \right]$$

$$\mathcal{I}(\underline{\theta}) = E \left[\frac{\partial \text{LNP}}{\partial \underline{\theta}} \frac{\partial \text{LNP}^T}{\partial \underline{\theta}} \right]$$

$$\begin{aligned}\underline{a}^T \mathcal{I}(\underline{\theta}) \underline{a} &= E \left[\underline{a}^T \frac{\partial \text{LNP}}{\partial \underline{\theta}} \frac{\partial \text{LNP}^T}{\partial \underline{\theta}} \underline{a} \right] \\ &= E \left[\left(\underline{a}^T \frac{\partial \text{LNP}}{\partial \underline{\theta}} \right)^2 \right] \geq 0\end{aligned}$$

for all $\underline{a} \Rightarrow \mathcal{I}(\underline{\theta})$ is positive semidefinite

for all $\underline{\theta}$.

From Prob 3.3 $\underline{\theta} = \begin{bmatrix} A \\ r \end{bmatrix}$ and using (3.31)

$$[\underline{I}(\underline{\theta})]_{ij} = \frac{1}{\sigma^2} \frac{\partial \underline{\mu}(\underline{\theta})}{\partial \theta_i} \frac{\partial \underline{\mu}(\underline{\theta})}{\partial \theta_j}$$

$$\text{where } \underline{\mu}(\underline{\theta}) = \begin{bmatrix} A \\ Ar \\ \vdots \\ Ar^{N-1} \end{bmatrix}$$

$$\frac{\partial \underline{\mu}(\underline{\theta})}{\partial A} = [1 \ r \ \dots \ r^{N-1}]^T$$

$$\frac{\partial \underline{\mu}(\underline{\theta})}{\partial r} = A [0 \ 1 \ \dots \ (N-1) \ r^{N-2}]^T$$

$$\underline{I}(\underline{\theta}) = \frac{1}{\sigma^2} \begin{bmatrix} \sum_{n=0}^{N-1} r^n & A \sum_{n=0}^{N-1} n r^{2n-1} \\ A \sum_{n=0}^{N-1} n r^{2n-1} & A^2 \sum_{n=0}^{N-1} n^2 r^{2n-2} \end{bmatrix}$$

If $A = 0$, $\underline{I}(\underline{\theta})$ is not positive definite since its determinant is zero. Clearly, in this case there is no information in the data about r .

- 11) Since $\underline{I}(\underline{\theta})$ is positive definite, $a > 0$, $c > 0$ and $\det(\underline{I}(\underline{\theta})) > 0$ or $ac - b^2 > 0$. But

$$(\mathbf{I}^{-1}(\underline{\theta}))_{ii} = \frac{c}{ac - b^2} = \frac{1}{a - b^2/c} \geq \frac{1}{a}$$

Thus, the CRLB is almost always increased when we estimate additional parameters.

Equality holds if and only if $b = 0$ or the Fisher information matrix is "decoupled", i.e., it is diagonal. In this case the additional parameter does not affect the CRLB.

$$12) \quad 1^2 = (\underline{e}_i^T \sqrt{\mathbf{I}(\underline{\theta})} \sqrt{\mathbf{I}^{-1}(\underline{\theta})} \underline{e}_i)^2$$

$$\text{Since } \sqrt{\mathbf{I}^{-1}(\underline{\theta})} = (\sqrt{\mathbf{I}(\underline{\theta})})^{-1}$$

$$1^2 \leq \underline{e}_i^T \mathbf{I}(\underline{\theta}) \underline{e}_i \quad \underline{e}_i^T \mathbf{I}^{-1}(\underline{\theta}) \underline{e}_i$$

$$1 \leq [\mathbf{I}(\underline{\theta})]_{ii} [\mathbf{I}^{-1}(\underline{\theta})]_{ii}$$

$$\Rightarrow [\mathbf{I}^{-1}(\underline{\theta})]_{ii} \geq \frac{1}{[\mathbf{I}(\underline{\theta})]_{ii}}$$

New bound achieved when an efficient estimator exists and $\mathbf{I}(\underline{\theta})$ is diagonal.

$$13) \text{ From (3.33) with } s(n; \underline{\theta}) = \sum_{k=0}^{p-1} A_k n^k$$

$$[I(\theta)]_{ij} = \frac{1}{\sigma^2} \sum_{n=0}^{N-1} \left(\frac{\partial}{\partial \theta_i} \sum_{k=0}^{p-1} A_k n^k \right) \left(\frac{\partial}{\partial \theta_j} \sum_{k=0}^{p-1} A_k n^k \right)$$

$$= \frac{1}{\sigma^2} \sum_{n=0}^{N-1} n^{i-1} n^{j-1} = \frac{1}{\sigma^2} \sum_{n=0}^{N-1} n^{i+j-2}$$

$$\text{where } \theta_i = A_{i-1} \quad i=1, 2, \dots, p$$

$$14) \quad \hat{A} = \frac{1}{N} \sum_{n=0}^{N-1} x(n)$$

If we condition the mean and variance on A , then we can regard A as the observed value A_0 . Hence, from Example 3.3

$$E(\hat{A} | A = A_0) = A_0$$

$$\text{var}(\hat{A} | A = A_0) = \sigma^2/N$$

and since $\sigma^2/N \rightarrow 0$ as $N \rightarrow \infty$, $\hat{A} \rightarrow A_0$.

Now consider A as a random variable

as $N \rightarrow \infty$

$$\begin{aligned} \text{var}(\hat{\sigma}_A^2) &= \text{var}(\hat{A}^2) \\ &\rightarrow \text{var}(A^2) \end{aligned}$$

$$\text{But } \text{var}(A^2) = E(A^4) - E(A^2)^2$$

$$= 3\sigma_A^4 - \sigma_A^4 = 2\sigma_A^4$$

so that $\text{var}(\hat{\sigma}_A^2) \rightarrow 2\sigma_A^4$, which is just the CRLB as $N \rightarrow \infty$. $\hat{\sigma}_A^2$ cannot be estimated without error since even as

$N \rightarrow \infty$, although we can nullify the noise effects by averaging ($\hat{A} \rightarrow A_0$), we cannot reduce the random nature of A . This is because we have only one realization of A . Since $\hat{\sigma}_A^2$ is the square of $\hat{A}$, it will also exhibit the same variability.

15) Because the $X(W)$'s are independent

$$I(p) = N i(p)$$

where $i(p)$ is the Fisher information for a single vector sample. Using (3.32)

$$i(p) = \frac{1}{2} \pi \left[\left(\underline{C}^{-1}(p), \frac{\partial \underline{C}(p)}{\partial p} \right)^2 \right]$$

$$\underline{C}^{-1}(p) = \frac{\begin{bmatrix} 1 & -p \\ -p & 1 \end{bmatrix}}{1-p^2} \quad \frac{\partial \underline{C}(p)}{\partial p} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$$

$$\underline{D} = \underline{C}^{-1}(p) \frac{\partial \underline{C}(p)}{\partial p} = \frac{\begin{bmatrix} -p & 1 \\ 1 & -p \end{bmatrix}}{1-p^2}$$

$$\underline{D}^2 = \frac{\begin{bmatrix} -p & 1 \\ 1 & -p \end{bmatrix} \begin{bmatrix} -p & 1 \\ 1 & -p \end{bmatrix}}{(1-p^2)^2} = \frac{\begin{bmatrix} 1+p^2 & - \\ - & 1+p^2 \end{bmatrix}}{(1-p^2)^2}$$

$$i(p) = \frac{1}{2} \frac{2+2p^2}{(1-p^2)^2} = \frac{1+p^2}{(1-p^2)^2}$$

$$\text{var}(\hat{\rho}) \geq \frac{(1-\rho^2)^2}{N(1+\rho^2)}$$

$$16) \quad \mathbf{I}(\rho_0) = \frac{1}{2} \text{tr} \left[\left(\underline{\mathbf{C}}^{-1}(\rho_0) \frac{\partial \underline{\mathbf{C}}(\rho_0)}{\partial \rho_0} \right)^2 \right]$$

$$\begin{aligned} \text{Let } r_{xx}[k] &= \mathcal{F}^{-1} \{ P_{xx}(f) \} \\ &= \mathcal{F}^{-1} \{ P_0 Q(f) \} \\ &= P_0 \mathcal{F}^{-1} \{ Q(f) \} \end{aligned}$$

Let $\mathcal{F}^{-1} \{ Q(f) \} = g[k]$ and construct the Toeplitz autocorrelation matrix $\underline{\mathbf{C}}_g$ (of dimension $N \times N$), where

$$(\underline{\mathbf{C}}_g)_{ij} = g[i-j]$$

Then, $\underline{\mathbf{C}}(\rho_0) = P_0 \underline{\mathbf{C}}_g$ and

$$\underline{\mathbf{C}}^{-1}(\rho_0) \frac{\partial \underline{\mathbf{C}}(\rho_0)}{\partial \rho_0} = \frac{1}{P_0} \underline{\mathbf{C}}_g^{-1} \underline{\mathbf{C}}_g = \frac{1}{P_0} \mathbf{I}$$

$$\mathbf{I}(\rho_0) = \frac{1}{2} \text{tr} \left[\frac{1}{P_0^2} \mathbf{I}^2 \right] = \frac{N}{2P_0^2}$$

$$\text{var}(\hat{\rho}_0) \geq 2P_0^2/N$$

Using the asymptotic form

$$\begin{aligned} \mathbf{I}(\rho_0) &= \frac{N}{2} \int_{-\frac{1}{2}}^{\frac{1}{2}} \left(\frac{\partial \ln P_{xx}(f; \rho_0)}{\partial \rho_0} \right)^2 df \\ &= \frac{N}{2} \int_{-\frac{1}{2}}^{\frac{1}{2}} \left(\frac{\partial \ln P_0 Q(f)}{\partial \rho_0} \right)^2 df \end{aligned}$$

$$= \frac{N}{2} \int_{-\frac{1}{2}}^{\frac{1}{2}} \frac{1}{P_0^2} df = \frac{N}{2P_0^2}$$

The two CRLB's are identical but in general this will not be true.

- 17) All elements of $\underline{I}(\underline{\theta})$ are the same except the sums run from $n = -M$ to $n = M$. Thus, now $(\underline{I}(\underline{\theta}))_{23} = 0$ since

$$\sum_{n=-M}^M n = 0$$

This makes $\underline{I}(\underline{\theta})$ diagonal. Letting $N = 2M + 1$

$$\text{var}(\hat{A}) \geq 2\sigma^2/N \quad \text{same as before}$$

$$\text{var}(\hat{\phi}) \geq \frac{2\sigma^2}{NA^2} = \frac{1}{N\eta} \quad \text{less than before}$$

$$\text{var}(\hat{f}_0) \geq \frac{\sigma^2}{2A^2\pi^2 \sum_{n=-M}^M n^2}$$

$$\text{But } \sum_{n=-M}^M n^2 = 2 \sum_{n=1}^M n^2 = \frac{2M(M+1)(2M+1)}{6}$$

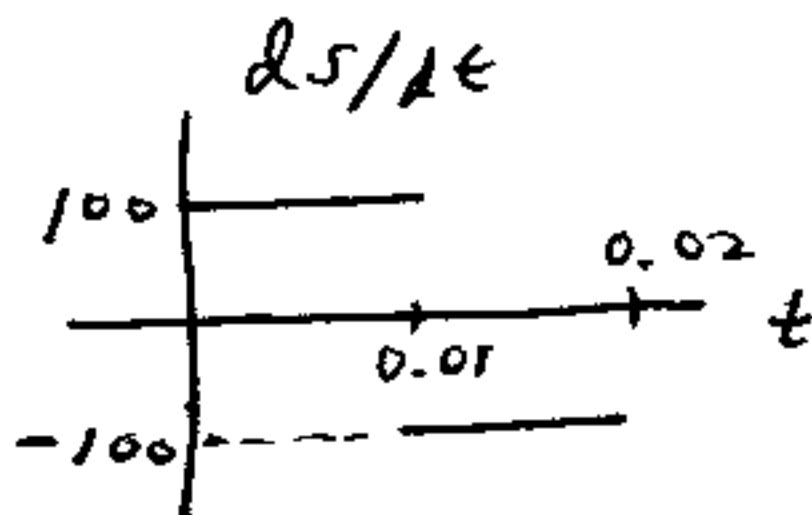
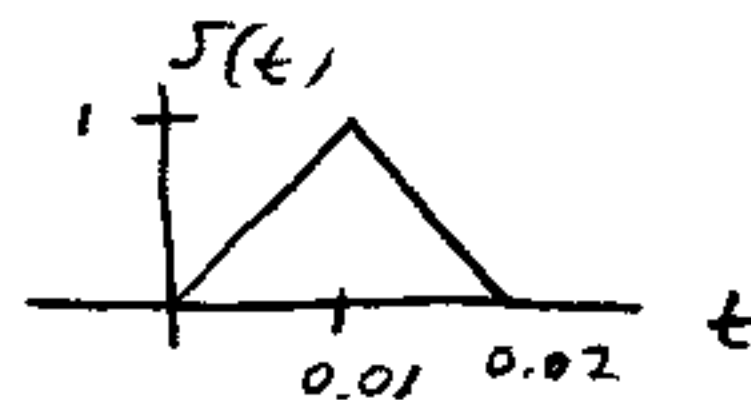
$$= \frac{1}{3} \left(\frac{N-1}{2} \right) \left(\frac{N+1}{2} \right) N = \frac{N(N^2-1)}{12}$$

$$\text{var}(\hat{f}_0) \geq \frac{6\sigma^2}{A^2\pi^2 N(N^2-1)} = \frac{3}{\eta\pi^2 N(N^2-1)}$$

$$= \frac{12}{(2\pi)^2 \eta N(N^2-1)} \quad \text{Same as before}$$

$$18) \quad \text{var}(\hat{R}) \geq \frac{C^2/4}{\frac{\epsilon}{N_0/2} \bar{F}^2}$$

$$\text{where } \bar{F}^2 = \frac{\int_0^{T_s} \left(\frac{ds}{dt} \right)^2 dt}{\int_0^{T_s} s^2(t) dt}$$



$$\bar{F}^2 = \frac{\int_0^{0.02} (100)^2 dt}{\epsilon} = \frac{200}{\epsilon}$$

$$\text{var}(\hat{R}) \geq \frac{C^2/4}{\frac{1}{N_0/2} 200} = \frac{(1500)^2/4}{10^6 \cdot 200}$$

$$= 0.00281$$

$$\text{or } \sqrt{\text{var}(\hat{R})} \geq 0.05 \text{ m}$$

19) From Example 3.15

$$\text{var}(\hat{\beta}) \geq \frac{12}{(2\pi)^2 M \eta \frac{M+1}{M-1} \left(\frac{L}{\lambda}\right)^2 \sin^2 \beta}$$

$$\text{For } \beta = 90^\circ, \eta = 1, F_0 = 10, L = (M-1)d \\ = (M-1)\lambda/2$$

$$\text{var}(\hat{\beta}) \geq \frac{12}{(2\pi)^2 M \frac{M+1}{M-1} \left(\frac{M-1}{2}\right)^2}$$

$$= \frac{12}{(2\pi)^2 \frac{M}{4} (M^2-1)}$$

$$M(M^2-1) \geq \frac{48}{(2\pi)^2 (5\pi/180)^2} = 159.7$$

or $M \geq 6$. But then

$$L = (M-1)\lambda/2 = (M-1) \frac{c}{2F_0} = \frac{5(3 \times 10^8)}{2 \times 10^6} \\ = 750 \text{ m}$$

This is clearly impossible.

$$\begin{aligned} 20) \quad \text{var}(\hat{P}_{xx}(f)) &\geq \frac{\left(\frac{\partial P_{xx}(f)}{\partial a(l)}\right)^2}{I(a(l))} \\ &= \frac{\left(\frac{\partial P_{xx}(f)}{\partial a(l)}\right)^2}{N/(1-a(l))^2} \end{aligned}$$

using results from Example 3.16.

$$P_{xx}(f) = \frac{\sigma_u^2}{|A(f)|^2} \quad \text{where } A(f) = 1 + a(1) e^{-j2\pi f}$$

$$\frac{\partial P_{xx}(f)}{\partial a(1)} = \sigma_u^2 \frac{\partial}{\partial a(1)} \left(\frac{1}{A(f)A^*(f)} \right)$$

$$= - \frac{\sigma_u^2}{|A(f)|^4} \frac{\partial}{\partial a(1)} A(f)A^*(f)$$

$$= - \frac{\sigma_u^2}{|A(f)|^4} (A(f) e^{j2\pi f} + A^*(f) e^{-j2\pi f})$$

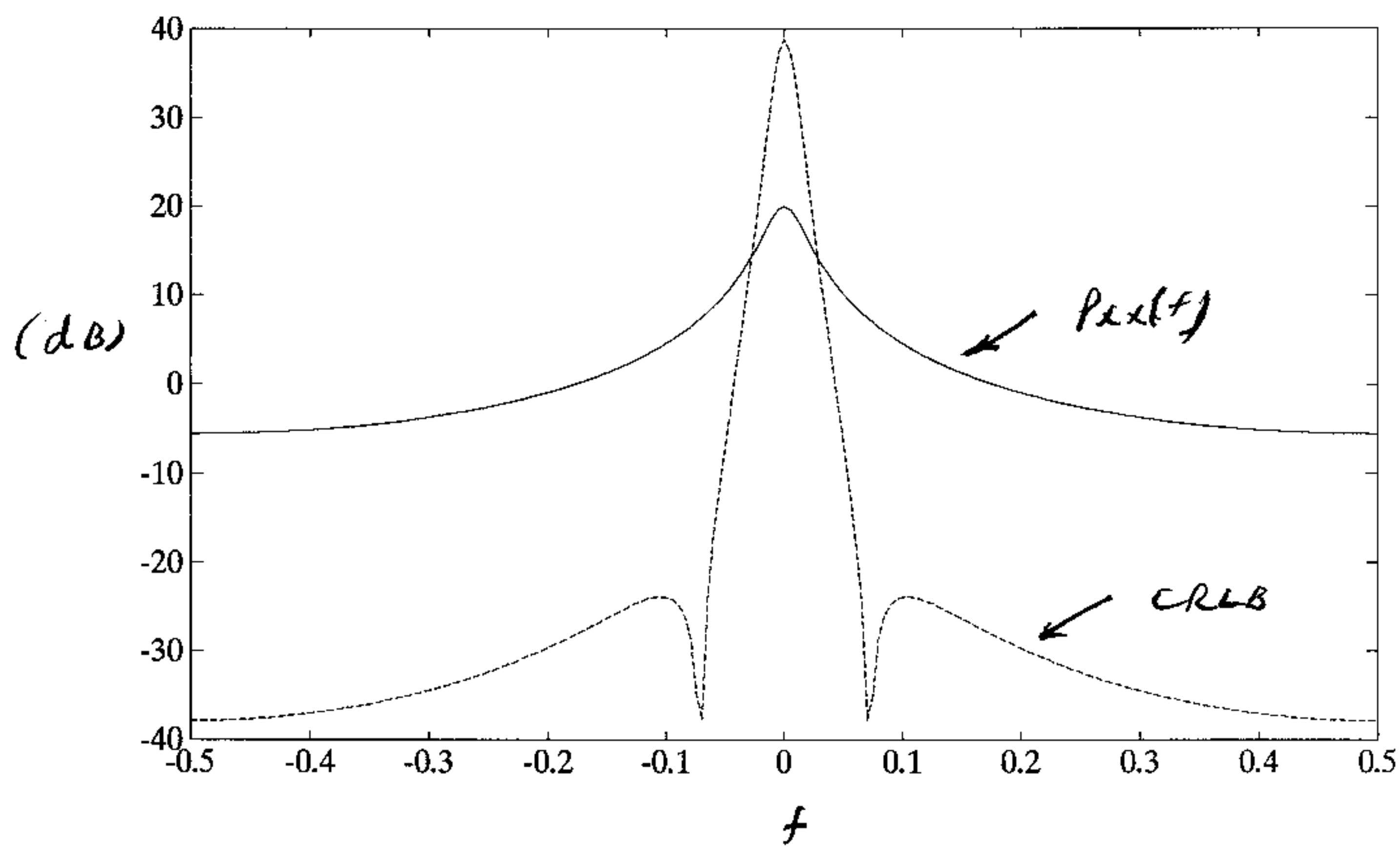
$$= \frac{-\sigma_u^2}{|A(f)|^4} \underbrace{2 \operatorname{Re}(A(f) e^{j2\pi f})}_{a(1) + \cos 2\pi f}$$

$$\operatorname{var}(\hat{P}_{xx}(f)) \geq \frac{\frac{4\sigma_u^4}{N}(1-a^2(1))}{|A(f)|^8} (a(1) + \cos 2\pi f)^2$$

For the given values

$$\begin{aligned} \operatorname{var}(\hat{P}_{xx}(f)) &\geq 0.0076 \frac{(a(1) + \cos 2\pi f)^2}{|A(f)|^8} \\ &= \frac{0.0076 (-0.9 + \cos 2\pi f)^2}{|1 - 0.9 e^{-j2\pi f}|^8} \end{aligned}$$

See Figure.



Prob. 3.20

Because of the sensitivity of the PSD to small changes in $a(\omega)$ for ω near zero, the variance is highest at DC. Note that

$$\begin{aligned} \left. \frac{\partial P_{xx}(\omega)}{\partial a(\omega)} \right|_{\omega=0} &= \frac{-\sigma_u^2 \cdot 2(1+a(\omega))}{(1+a(\omega))^4} \\ &= \frac{-2\sigma_u^2}{(1+a(\omega))^3} \end{aligned}$$

and for this example

$$\left[\left. \frac{\partial P_{xx}(\omega)}{\partial a(\omega)} \right|_{\omega=0} \right]^2 = 4 \times 10^6$$

Chapter 4

1) This fits linear model form.

$$\underline{X} = \underline{H}\underline{\theta} + \underline{W}$$

$$\underline{H} = \begin{bmatrix} 1 & 1 & \dots & 1 \\ r_1 & r_2 & \dots & r_p \\ \vdots & \vdots & & \vdots \\ r_1^{N-1} & r_2^{N-1} & \dots & r_p^{N-1} \end{bmatrix} \quad \underline{\theta} = \begin{bmatrix} A_1 \\ A_2 \\ \vdots \\ A_p \end{bmatrix}$$

$$\hat{\underline{\theta}} = (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{X} \quad C_{\hat{\underline{\theta}}} = \sigma^2 (\underline{H}^T \underline{H})^{-1}$$

For $p=2$, $r_1=1$, $r_2=-1$ and N even

$$\underline{H} = \begin{bmatrix} 1 & 1 \\ 1 & -1 \\ \vdots & \vdots \\ 1 & -1 \end{bmatrix} \Rightarrow \underline{H}^T \underline{H} = \begin{bmatrix} N & 0 \\ 0 & N \end{bmatrix} = N \underline{I}$$

since columns are orthogonal

$$\hat{\underline{\theta}} = \frac{1}{N} \underline{H}^T \underline{X} = \begin{bmatrix} \frac{1}{N} \sum_{n=0}^{N-1} X[n] \\ \frac{1}{N} \sum_{n=0}^{N-1} (-1)^n X[n] \end{bmatrix}$$

$$C_{\hat{\underline{\theta}}} = \frac{\sigma^2}{N} \underline{I}$$

2) First assume columns of $\underline{H}$ are linearly independent
Then, $\underline{H}\underline{x} = \sum_{i=1}^p x_i \underline{h}_i \neq 0$ if $\underline{x} \neq 0$

$$\Rightarrow \underline{x}^T \underline{H}^T \underline{H} \underline{x} = \|\underline{H}\underline{x}\|^2 > 0 \text{ for } \underline{x} \neq 0$$

$\Rightarrow \underline{H}^T \underline{H}$ is positive definite

Now assume $\underline{H}^T \underline{H}$ is positive definite or

$$\underline{x}^T \underline{H}^T \underline{H} \underline{x} > 0 \quad \text{for all } \underline{x} \neq \underline{0}$$

$$\text{or} \quad \|\underline{H} \underline{x}\|^2 > 0 \quad \text{for all } \underline{x} \neq \underline{0}$$

$$\Rightarrow \underline{H} \underline{x} \neq \underline{0} \quad \text{for all } \underline{x} \neq \underline{0}$$

$\Rightarrow$ columns of $\underline{H}$ are linearly independent

It can further be shown that for matrices of the form $\underline{H}^T \underline{H}$, invertibility is equivalent to being positive definite.

$$3) \quad \underline{H}^T \underline{H} = \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1+\epsilon \end{bmatrix} \begin{bmatrix} 1 & 1 \\ 1 & 1 \\ 1 & 1+\epsilon \end{bmatrix} = \begin{bmatrix} 3 & 3+\epsilon \\ 3+\epsilon & 2+(1+\epsilon)^2 \end{bmatrix}$$

$$(\underline{H}^T \underline{H})^{-1} = \frac{1}{\begin{vmatrix} 2+(1+\epsilon)^2 & -(3+\epsilon) \\ -(3+\epsilon) & 3 \end{vmatrix}} \begin{bmatrix} 2+(1+\epsilon)^2 & -(3+\epsilon) \\ -(3+\epsilon) & 3 \end{bmatrix}$$

$$= \frac{1}{3[2+(1+\epsilon)^2] - (3+\epsilon)^2}$$

$$= \frac{1}{2\epsilon^2} \begin{bmatrix} 3+2\epsilon+\epsilon^2 & -(3+\epsilon) \\ -(3+\epsilon) & 3 \end{bmatrix}$$

$$2\epsilon^2$$

As $\epsilon \rightarrow 0$, all elements $\rightarrow \infty$

$$\begin{aligned}
\hat{\underline{\theta}} &= (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{x} \\
&= \frac{1}{2\epsilon^2} \begin{bmatrix} 3+2\epsilon+\epsilon^2 & -(3+\epsilon) \\ -(3+\epsilon) & 3 \end{bmatrix} \begin{bmatrix} 6 \\ 6+2\epsilon \end{bmatrix} \\
&= \frac{1}{2\epsilon^2} \begin{bmatrix} 18+12\epsilon+6\epsilon^2-18-6\epsilon-6\epsilon-2\epsilon^2 \\ -18-6\epsilon+18+6\epsilon \end{bmatrix} \\
&= \begin{bmatrix} 2 \\ 0 \end{bmatrix}
\end{aligned}$$

Hence, even as $\epsilon \rightarrow 0$, $\hat{\underline{\theta}} = \begin{bmatrix} 2 \\ 0 \end{bmatrix}$. This is because $\underline{x}$ lies in the subspace spanned by the first column of $\underline{H}$, which does not depend on ϵ .

$$4) \quad \hat{\underline{\theta}} \sim N(\underline{\theta}, \sigma^2 (\underline{H}^T \underline{H})^{-1})$$

$$\hat{\underline{y}} = \underline{H} \hat{\underline{\theta}} \sim N(\underline{H} \underline{\theta}, \sigma^2 \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T)$$

Note that the covariance matrix is singular since $\underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T$ is a projection matrix of rank p . (See Chapter 8).

For Example 4.2

$$\hat{y}(n) = \sum_{k=1}^M \hat{a}_k \cos \frac{2\pi k n}{N} + \sum_{k=1}^M \hat{b}_k \sin \frac{2\pi k n}{N}$$

where $\hat{a}_k = \frac{2}{N} \sum_{n=0}^{N-1} x(n) \cos 2\pi k \frac{n}{N}$

$$\hat{b}_k = \frac{2}{N} \sum_{n=0}^{N-1} x(n) \sin 2\pi k \frac{n}{N}$$

Also, $(\underline{H}^T \underline{H})^{-1} = \frac{2}{N} \underline{I} \Rightarrow$

$$\underline{\hat{y}} \sim N(\underline{y}, \frac{2\sigma^2}{N} \underline{H} \underline{H}^T) \quad \text{where } \underline{y} = \underline{H} \underline{\theta}.$$

5) Let $\omega_k = \frac{2\pi}{N} k n$

$$\sum_{n=0}^{N-1} \cos \omega_k \cos \omega_l = \frac{1}{2} \sum_n [\cos(\omega_k + \omega_l) + \cos(\omega_k - \omega_l)]$$

$$= \frac{1}{2} \operatorname{Re} \sum_n e^{j(\omega_k + \omega_l)} + \frac{1}{2} \operatorname{Re} \sum_n e^{j(\omega_k - \omega_l)}$$

But $\sum_{n=0}^{N-1} e^{j(\omega_k + \omega_l)} = \sum_{n=0}^{N-1} e^{j \frac{2\pi}{N} (k+l)n}$

$$= \frac{1 - e^{j \frac{2\pi}{N} (k+l)N}}{1 - e^{j \frac{2\pi}{N} (k+l)}} = \frac{1 - e^{j 2\pi (k+l)}}{1 - e^{j \frac{2\pi}{N} (k+l)}} = 0$$

And similarly for $\sum_n e^{j(\omega_k - \omega_l)}$.

b) From Example 4.2

$$\hat{a}_k \sim N(a_k, 2\sigma^2/N) \quad \hat{b}_k \sim N(b_k, 2\sigma^2/N)$$

and $\hat{a}_k, \hat{b}_k$ are independent

$$\begin{aligned}
 E(\hat{P}) &= \frac{E(\hat{a}_k^2) + E(\hat{b}_k^2)}{2} \\
 &= \frac{\text{var}(\hat{a}_k) + E(\hat{a}_k)^2 + \text{var}(\hat{b}_k) + E(\hat{b}_k)^2}{2} \\
 &= \frac{a_k^2 + b_k^2 + 4\sigma^2/N}{2} = P + \frac{2\sigma^2}{N}
 \end{aligned}$$

$$\text{var}(\hat{P}) = \frac{\text{var}(\hat{a}_k^2) + \text{var}(\hat{b}_k^2)}{4}$$

since $\hat{a}_k, \hat{b}_k$ are independent

But $\text{var}(\hat{a}_k^2)$ can be found as follows:

$$\text{If } Z \sim N(\mu, \sigma^2)$$

$$\text{var}(Z^2) = 4\mu^2\sigma^2 + 2\sigma^4$$

(see development preceding (3.19))

$$\Rightarrow \text{var}(\hat{a}_k^2) = 4a_k^2 \frac{2\sigma^2}{N} + 2\left(\frac{2\sigma^2}{N}\right)^2$$

$$\text{var}(\hat{b}_k^2) = 4b_k^2 \frac{2\sigma^2}{N} + 2\left(\frac{2\sigma^2}{N}\right)^2$$

$$\text{var}(\hat{P}) = 2P \frac{2\sigma^2}{N} + \left(\frac{2\sigma^2}{N}\right)^2 = \frac{2\sigma^2}{N} \left(2P + \frac{2\sigma^2}{N}\right)$$

$$\frac{E(\hat{P})^2}{\text{var}(\hat{P})} = \frac{(P + 2\sigma^2/N)^2}{\frac{2\sigma^2}{N} (2P + \frac{2\sigma^2}{N})}$$

When no signal is present or $P = 0$, the measure is one.

For example, if $p \gg 4\sigma^2/N$ or $NP \gg 4\sigma^2$,

$$\frac{E(\hat{P})^2}{\text{var}(\hat{P})} = \frac{p^2}{4p\sigma^2/N} = \frac{p}{\frac{4\sigma^2}{N}} \gg 1$$

and the sinusoid will be easily detectable.

$$7) \quad [\underline{H}^T \underline{H}]_{ij} = \sum_{n=1}^N u[n-i]u[n-j]$$

For large N this is $\sum_{n=-\infty}^{\infty} u[n-i]u[n-j]$ since for $u[n]=0$, $n \leq 0$ and $n \geq N-1$, this will add only about $2p$ additional terms. If $N \gg p$, these terms will be negligible.

Now assume $i \geq j$ and let $m = n - i$

$$\begin{aligned} [\underline{H}^T \underline{H}]_{ij} &= \sum_{m=-\infty}^{\infty} u[m]u[m+i-j] \\ &= \sum_{m=0}^{N-1-(i-j)} u[m]u[m+i-j] \end{aligned}$$

Similarly for $i < j$

$$[\underline{H}^T \underline{H}]_{ij} = \sum_{m=0}^{N-1-(j-i)} u[m]u[m+j-i]$$

$$\Rightarrow [\underline{H}^T \underline{H}]_{ij} = \sum_{m=0}^{N-1-|i-j|} u[m]u[m+|i-j|]$$

$$8) \quad x[n] = \sum_{l=0}^{\infty} h[l] u[n-l]$$

$$\begin{aligned} r_{ux}[k] &= E \left[u[n] x[n+k] \right] \\ &= E \left[u[n] \sum_{l=0}^{\infty} h[l] u[n+k-l] \right] \\ &= \sum_{l=0}^{\infty} h[l] E \left[u[n] u[n+k-l] \right] \\ &= \sum_{l=0}^{\infty} h[l] r_{uu}[k-l] \\ &= h[k] \star r_{uu}[k] \end{aligned}$$

$$\Rightarrow P_{ux}(f) = H(f) P_{uu}(f)$$

If $P_{uu}(f) = \sigma^2$, then $H(f) = P_{ux}(f) / \sigma^2$
or

$$\hat{H}(f) = \frac{\hat{P}_{ux}(f)}{r_{uu}[0]}$$

If $P_{uu}(f) = 0$ over a band of frequencies, it would be impossible to estimate $H(f)$ over that band. This is because there would be no power over that band at the output, leading to $\hat{P}_{ux}(f) = 0$, independent of $H(f)$. The TDL estimator of (4.23) is just the inverse Fourier transform of $\hat{H}(f)$.

$$9) \quad p(\underline{x}; \underline{\theta}) = \frac{1}{(2\pi)^{N/2} \det(\underline{C})^{1/2}} e^{-\frac{1}{2}(\underline{x} - \underline{H}\underline{\theta})^T \underline{C}^{-1}(\underline{x} - \underline{H}\underline{\theta})}$$

$$\begin{aligned} \frac{\partial \ln p}{\partial \underline{\theta}} &= -\frac{1}{2} \frac{\partial}{\partial \underline{\theta}} \left[\underline{x}^T \underline{C}^{-1} \underline{x} - 2 \underline{\theta}^T \underline{H}^T \underline{C}^{-1} \underline{x} + \underline{\theta}^T \underline{H}^T \underline{C}^{-1} \underline{H} \underline{\theta} \right] \\ &= -\frac{1}{2} \left[-2 \underline{H}^T \underline{C}^{-1} \underline{x} + 2 \underline{H}^T \underline{C}^{-1} \underline{H} \underline{\theta} \right] \text{ using (4.3)} \\ &= \underline{H}^T \underline{C}^{-1} \underline{x} - \underline{H}^T \underline{C}^{-1} \underline{H} \underline{\theta} \\ &= \underbrace{\underline{H}^T \underline{C}^{-1} \underline{H}}_{\underline{I}(\underline{\theta})} \underbrace{(\underline{H}^T \underline{C}^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}^{-1} \underline{x} - \underline{\theta}}_{\hat{\underline{\theta}}} \end{aligned}$$

$\therefore \hat{\underline{\theta}} = (\underline{H}^T \underline{C}^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}^{-1} \underline{x}$ is MVU estimator (and efficient)

$$\underline{C}_{\hat{\underline{\theta}}} = (\underline{H}^T \underline{C}^{-1} \underline{H})^{-1}$$

$$10) \quad \underline{C}^{-1} = \underline{D}^T \underline{D}$$

$$\text{Since } \underline{C} = \text{diag}(\sigma_0^2, \sigma_1^2, \dots, \sigma_{N-1}^2)$$

$$\underline{C}^{-1} = \text{diag}(1/\sigma_0^2, 1/\sigma_1^2, \dots, 1/\sigma_{N-1}^2)$$

$$\Rightarrow \underline{D} = \text{diag}(1/\sigma_0, 1/\sigma_1, \dots, 1/\sigma_{N-1})$$

$$\hat{A} = \sum_{n=0}^{N-1} d_n x[n]$$

$$\text{where } d_n = \frac{[\underline{D}^{-1}]_n}{\underline{1}^T \underline{D}^{-1} \underline{1}} = \frac{1/\sigma_n}{\sum_{m=0}^{N-1} 1/\sigma_m^2}$$

Since $\underline{C}$ is already diagonal or the

components of $\underline{w}$ are uncorrelated, we need only form $\underline{x}' = \underline{D}\underline{x}$ or $x'(n) = x(n)/\sigma_n$ so that all variances are one. Then, we "average" the $x'(n)$ samples. Actually, we weight the samples since the DC level has been changed to a non-DC signal due to the prewhitening stage.

If a $\sigma_n^2 = 0$, say σ_m^2 , then we cannot prewhiten the data. In this case, however, as $\sigma_m^2 \rightarrow 0$, $d_n \rightarrow 0$ for $n \neq m$ and $d_n \rightarrow \sigma_m$ for $n = m$. Thus, $\hat{A} \rightarrow d_m x'(m) = x(m)$, as expected.

$$11) \quad \hat{A} = (\underline{H}^T \underline{C}^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}^{-1} \underline{x} \quad \text{var}(\hat{A}) = (\underline{H}^T \underline{C}^{-1} \underline{H})^{-1}$$

$$\underline{C} = \sigma^2 \begin{pmatrix} 1 & p \\ p & 1 \end{pmatrix} \quad \underline{H} = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$$

$$\underline{C}^{-1} = \frac{1}{\sigma^2(1-p^2)} \begin{pmatrix} 1 & -p \\ -p & 1 \end{pmatrix}$$

$$\underline{H}^T \underline{C}^{-1} \underline{H} = \frac{2-2p}{\sigma^2(1-p^2)} = \frac{2(1-p)}{\sigma^2(1-p^2)} = \frac{2}{\sigma^2(1+p)}$$

$$\begin{aligned} \underline{H}^T \underline{C}^{-1} \underline{x} &= \frac{1^T}{\sigma^2(1-p^2)} \begin{pmatrix} 1 & -p \\ -p & 1 \end{pmatrix} \begin{pmatrix} x(0) \\ x(1) \end{pmatrix} \\ &= \frac{x(0)(1-p) + x(1)(1-p)}{\sigma^2(1-p^2)} \end{aligned}$$

$$= \frac{x(0) + x(1)}{\sigma^2(1+p)}$$

$$\hat{A} = \frac{\sigma^2(1+p)}{2} \frac{x(0) + x(1)}{\sigma^2(1+p)} = \frac{1}{2}(x(0) + x(1))$$

$$\text{var}(\hat{A}) = \frac{\sigma^2(1+p)}{2}$$

We don't need prewhitening here because $\underline{H}$ is an eigenvector of $\underline{C}$. Hence,

$$\begin{aligned} (\underline{H}^T \underline{C}^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}^{-1} &= (\underline{H}^T \frac{1}{\lambda} \underline{H})^{-1} \frac{1}{\lambda} \underline{H}^T \\ &= (\underline{H}^T \underline{H})^{-1} \underline{H}^T \end{aligned}$$

As $\rho \rightarrow 1$, $\text{var}(\hat{A}) \rightarrow \sigma^2$

As $\rho \rightarrow -1$, $\text{var}(\hat{A}) \rightarrow 0$.

See Prob 3.9 for explanation.

$$12) \quad \text{If } \underline{x} = \underline{H}\underline{\theta} + \underline{N}, \quad \underline{x}' = \underline{A}(\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{x}$$

$$\begin{aligned} \underline{x}' &= \underline{A}(\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{H} \underline{\theta} + \underline{A}(\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{N} \\ \uparrow &= \underline{A} \underline{\theta} + \underline{N}' = \underline{H}' \underline{\alpha} + \underline{N}' \quad \text{where } \underline{H}' = \underline{I} \\ r \times 1 & \quad \quad \quad \uparrow \\ & \quad \quad \quad r \times 1 \end{aligned}$$

$$\begin{aligned} \underline{C}' &= E(\underline{N}' \underline{N}'^T) = E(\underline{A}(\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{N} \underline{N}^T \underline{H}(\underline{H}^T \underline{H})^{-1} \underline{A}^T) \\ &= \sigma^2 \underline{A}(\underline{H}^T \underline{H})^{-1} \underline{A}^T \end{aligned}$$

Since A is full rank, C^{-1} is positive definite and C^{-1} exists.

$$\Rightarrow \hat{\underline{\alpha}} = (\underline{H}^T \underline{C}^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}^{-1} \underline{x}$$

$$= \underline{C}^{-1} \underline{C}^{-1} \underline{x} = \underline{A} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{x} = \underline{A} \hat{\underline{\theta}}$$

$$\begin{aligned} 13) \quad E(\hat{\underline{\theta}}) &= E\left\{ (\underline{H}^T \underline{H})^{-1} \underline{H}^T (\underline{H} \underline{\theta} + \underline{w}) \right\} \\ &= \underline{\theta} + E\left\{ (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{w} \right\} \\ &= \underline{\theta} + E\left\{ (\underline{H}^T \underline{H})^{-1} \underline{H}^T \right\} E(\underline{w}) \\ &= \underline{\theta} \quad \text{since } E(\underline{w}) = \underline{0}. \end{aligned}$$

$$\begin{aligned} \underline{C}_{\hat{\theta}} &= E\left\{ (\hat{\underline{\theta}} - \underline{\theta})(\hat{\underline{\theta}} - \underline{\theta})^T \right\} \\ &= E\left\{ \left((\underline{H}^T \underline{H})^{-1} \underline{H}^T (\underline{x} - \underline{H} \underline{\theta}) \right) \left((\underline{H}^T \underline{H})^{-1} \underline{H}^T (\underline{x} - \underline{H} \underline{\theta}) \right)^T \right\} \\ &= E\left\{ (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{w} \underline{w}^T \underline{H} (\underline{H}^T \underline{H})^{-1} \right\} \\ &= E_{H|W} E_W \left\{ (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{w} \underline{w}^T \underline{H} (\underline{H}^T \underline{H})^{-1} \right\} \\ &= E_{H|W} \left[(\underline{H}^T \underline{H})^{-1} \underline{H}^T \sigma^2 \underline{I} \underline{H} (\underline{H}^T \underline{H})^{-1} \right] \\ &= E_{H|W} \left[\sigma^2 (\underline{H}^T \underline{H})^{-1} \right] = \sigma^2 E_H \left[(\underline{H}^T \underline{H})^{-1} \right] \end{aligned}$$

Since $\underline{H}$ and $\underline{w}$ are independent. If $\underline{H}$ and $\underline{w}$ are not independent $\hat{\underline{\theta}}$ may be biased.

$$14) \quad \underline{H} = \underline{1} \quad \text{with probability } 1 - \epsilon$$

$$\underline{H} = \left[\underbrace{11 \dots 1}_M \underbrace{00 \dots 0}_{N-M} \right]^T \quad \text{with probability } \epsilon$$

$$\hat{\underline{A}} = (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{x}$$

$$= \frac{1}{N} \sum_{n=0}^{N-1} x[n] \quad \text{no fade}$$

$$\frac{1}{M} \sum_{n=0}^{M-1} x[n] \quad \text{fade}$$

$$\text{var}(\hat{\underline{A}}) = \sigma^2 E_H [(\underline{H}^T \underline{H})^{-1}]$$

$$= \sigma^2 \left[\frac{1}{N} (1 - \epsilon) + \frac{1}{M} \epsilon \right]$$

$$= \frac{\sigma^2}{N} \left[1 - \epsilon + \frac{N}{M} \epsilon \right]$$

$$= \frac{\sigma^2}{N} \left[1 + \left(\frac{N}{M} - 1 \right) \epsilon \right] > \sigma^2/N$$

Clearly, the variances are the same only if $M = N$ or $\epsilon = 0$. Otherwise, it is increased.

Chapter 5

$$\begin{aligned}
 1) \quad p(\underline{x} \mid T(\underline{x}) = T_0; \sigma^2) &= \frac{p(\underline{x}; \sigma^2) \delta(T(\underline{x}) - T_0)}{p(T(\underline{x}) = T_0; \sigma^2)} \\
 &= \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_n x^2(n)} \delta(T(\underline{x}) - T_0) \\
 &\quad p\left(\sum_n x^2(n) = T_0; \sigma^2\right)
 \end{aligned}$$

But $\sum_n x^2(n) = T_0$

$$p\left(\sum_n x^2(n)\right) = \frac{1}{\sigma^2} p_r\left(\sum_n x^2(n) \mid \sigma^2\right)$$

$$\begin{aligned}
 p(\underline{x} \mid T(\underline{x}) = T_0; \sigma^2) &= \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} T_0} \delta(T(\underline{x}) - T_0) \\
 &\quad \frac{1}{\sigma^2} \frac{1}{2^{N/2} \Gamma(N/2)} e^{-T_0/2\sigma^2} \left(\frac{T_0}{\sigma^2}\right)^{N/2-1} \\
 &= \frac{1}{\pi^{N/2}} \delta(T(\underline{x}) - T_0) \\
 &\quad \frac{1}{\Gamma(N/2)} T_0^{N/2-1}
 \end{aligned}$$

$$\begin{aligned}
 2) \quad p(\underline{x}; \sigma^2) &= \prod_{n=0}^{N-1} \frac{x(n)}{\sigma^2} e^{-\frac{1}{2} x^2(n)/\sigma^2} \quad \text{all } x(n) > 0 \\
 &= \underbrace{u(\min x(n)) \prod_{n=0}^{N-1} x(n)}_{h(\underline{x})} \underbrace{\frac{1}{\sigma^{2N}} e^{-\frac{1}{2} \sum_n x^2(n)/\sigma^2}}_{g(T(\underline{x}), \sigma^2)}
 \end{aligned}$$

where $u(x)$ is the unit step function

$T(\underline{x}) = \sum_{n=0}^{N-1} x^2[n]$ is a sufficient statistic

$$\begin{aligned} 3) \quad p(\underline{x}; \lambda) &= \prod_{n=0}^{N-1} \lambda e^{-\lambda x[n]} \quad \text{all } x[n] > 0 \\ &= \underbrace{\lambda^N e^{-\lambda \sum_n x[n]}}_{g(T(\underline{x}), \lambda)} \cdot \underbrace{u(\min x[n])}_{h(\underline{x})} \end{aligned}$$

$T(\underline{x}) = \sum_{n=0}^{N-1} x[n]$ is a sufficient statistic

$$\begin{aligned} 4) \quad p(\underline{x}; \theta) &= \prod_{n=0}^{N-1} \frac{1}{\sqrt{2\pi\theta}} e^{-\frac{1}{2\theta}(x[n]-\theta)^2} \\ &= \frac{1}{(2\pi\theta)^{N/2}} e^{-\frac{1}{2\theta} \sum_n (x[n]-\theta)^2} \end{aligned}$$

But $\sum_n (x[n]-\theta)^2 = \sum_n x^2[n] - 2\theta \sum_n x[n] + N\theta^2$

$$p(\underline{x}; \theta) = \underbrace{\frac{1}{(2\pi\theta)^{N/2}} e^{-\frac{1}{2\theta} \sum_n x^2[n] - \frac{1}{2} N\theta}}_{g(T(\underline{x}), \theta)} \underbrace{e^{\sum_n x[n]}}_{h(\underline{x})}$$

$T(\underline{x}) = \sum_{n=0}^{N-1} x^2[n]$ is a sufficient statistic

$$\begin{aligned} 5) \quad p(x[n]; \theta) &= \frac{1}{2\theta} (u(x[n]+\theta) - u(x[n]-\theta)) \\ p(\underline{x}; \theta) &= \frac{1}{(2\theta)^N} \prod_{n=0}^{N-1} [u(x[n]+\theta) - u(x[n]-\theta)] \end{aligned}$$

But the product is zero unless $-\theta \leq x[n] \leq \theta$
for all $x[n]$ or $\min x[n] \geq -\theta$, $\max x[n] \leq \theta$

or $\max |x[n]| \leq \theta$ so that

$$p(\underline{x}; \theta) = \underbrace{\frac{1}{(2\theta)^N} u(\theta - \max |x[n]|)}_{g(T(\underline{x}), \theta)} \cdot \underbrace{1}_{h(\underline{x})}$$

and $T(\underline{x}) = \max |x[n]|$ is the sufficient statistic.

$$b) \quad p(\underline{x}; \sigma^2) = \underbrace{\frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_{n=0}^{N-1} (x[n] - A)^2}}_{g(T(\underline{x}), \sigma^2)} \cdot \underbrace{1}_{h(\underline{x})}$$

$T(\underline{x}) = \sum_{n=0}^{N-1} (x[n] - A)^2$ is a sufficient statistic

To make it unbiased, divide by N . Thus,

$\hat{\sigma}^2 = \frac{1}{N} \sum_{n=0}^{N-1} (x[n] - A)^2$ is the MVU estimator.

$$7) \quad p(\underline{x}; f_0) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_{n=0}^{N-1} (x[n] - \cos 2\pi f_0 n)^2}$$

$$= \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \left[\sum_n x^2[n] - 2 \sum_n x[n] \cos 2\pi f_0 n + \sum_n \cos^2 2\pi f_0 n \right]}$$

Because of the $\sum_n x[n] \cos 2\pi f_0 n$ term, which cannot be separated into a statistic, there does not appear to be a sufficient statistic.

Note that $\sum_n x[n] \cos 2\pi f_0 n$ is not a statistic

since it depends on f_0 .

$$8) \quad p(\underline{x}; r) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_{n=0}^{N-1} (x[n] - r^n)^2}$$

$$\text{But } \sum_n (x[n] - r^n)^2 = \sum_n x^2[n] - 2 \sum_n x[n] r^n + \sum_n r^{2n}$$

Again, the term $\sum_n x[n] r^n$ cannot be separated into a single sufficient statistic and a function of r .

$$9) \quad p_r \{x[n]\} = \theta^{x[n]} (1-\theta)^{1-x[n]} \quad x[n] = 0, 1$$

$$\begin{aligned} p_r \{ \underline{x} \} &= \prod_{n=0}^{N-1} \theta^{x[n]} (1-\theta)^{1-x[n]} \\ &= \underbrace{\theta^{\sum_n x[n]} (1-\theta)^{N - \sum_n x[n]}}_{g(T(\underline{x}), \theta)} \cdot \underbrace{1}_{h(\underline{x})} \end{aligned}$$

where $T(\underline{x}) = \sum_{n=0}^{N-1} x[n]$ is a sufficient statistic. To make it unbiased divide by N so that

$$\hat{\theta} = \frac{1}{N} \sum_{n=0}^{N-1} x[n] \text{ is the MVU estimator.}$$

10) At high SNR we ignore the noise so that

$$T_1(\underline{x}) \approx \sum_n A \cos(2\pi f_0 n + \phi) \cos 2\pi f_0 n$$

$$= \sum_n \frac{A}{2} [\cos \phi + \cos(4\pi f_0 n + \phi)]$$

$$\hat{\phi} \approx NA/2 \cos \phi \quad \text{using results of Prob 3.7}$$

Also,

$$T_2(\underline{x}) \approx \sum_n A \cos(2\pi f_0 n + \phi) \sin 2\pi f_0 n$$

$$= \sum_n \frac{A}{2} [-\sin \phi + \sin(4\pi f_0 n + \phi)]$$

$$\hat{\phi} \approx -\frac{NA}{2} \sin \phi$$

$$\text{Thus, } \hat{\phi} = -\arctan \frac{T_2(\underline{x})}{T_1(\underline{x})} = -\arctan \frac{-\frac{NA}{2} \sin \phi}{\frac{NA}{2} \cos \phi} \\ = \phi$$

This is not the MVU estimator since it is not unbiased. To see this note that even if we could assume $E(T_1(\underline{x})) = \frac{NA}{2} \cos \phi$ and

$$E(T_2(\underline{x})) = -\frac{NA}{2} \sin \phi$$

(which will not be true in general - only at high SNR), it is not true that

$$E(\hat{\phi}) = E\left[-\arctan \frac{T_2(\underline{x})}{T_1(\underline{x})}\right] \\ = -\arctan \frac{E\{T_2(\underline{x})\}}{E\{T_1(\underline{x})\}} \\ \uparrow \\ \text{incorrect}$$

$$11) \quad \theta = 2A + 1 \Rightarrow A = \frac{\theta-1}{2}$$

$$p(\underline{x}; A) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_n (x(n) - A)^2}$$

$$p'(\underline{x}; \theta) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_n (x(n) - \frac{\theta-1}{2})^2}$$

$$= \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \left[\sum_n x^2(n) - (\theta-1) \sum_n x(n) + N \left(\frac{\theta-1}{2} \right)^2 \right]}$$

$$= \underbrace{\frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \left[N \left(\frac{\theta-1}{2} \right)^2 - (\theta-1) \sum_n x(n) \right]}}_{g(T(\underline{x}), \theta)} \underbrace{e^{-\frac{1}{2\sigma^2} \sum_n x^2(n)}}_{h(\underline{x})}$$

Thus, $\sum_{n=0}^{N-1} x(n)$ is a sufficient statistic for θ .
To make it unbiased

$$E\left(\sum_n x(n)\right) = NA = N\left(\frac{\theta-1}{2}\right) = N\frac{\theta}{2} - \frac{N}{2}$$

$$\text{Let } \hat{\theta} = \frac{2}{N} \left(\sum_n x(n) + \frac{N}{2} \right) = \frac{2}{N} \sum_n x(n) + 1$$

This is the MVU estimator. Note that $\hat{\theta} = 2\hat{A} + 1$, where $\hat{A}$ is MVU estimator for A .

For $\theta = A^3$ the sufficient statistic is again $\sum x(n)$. To make it unbiased we need a function g such that

$$E\left(g\left(\sum_n x(n)\right)\right) = A^3$$

$$\text{or } E(h(\bar{X})) = A^3 \quad \text{where } \bar{X} \sim N(A, \sigma^2/N)$$

since $\bar{X}$ is also a sufficient statistic for θ .

Examining $\bar{X}^3$ we have

$$E[(\bar{X} - A)^3] = E(\bar{X}^3) - 3AE(\bar{X}^2) + 3A^2E(\bar{X}) - A^3 = 0$$

since all odd-order moments are zero.

$$E(\bar{X}^3) = 3A(A^2 + \sigma^2/N) - 3A^3 + A^3 = A^3 + 3A\sigma^2/N$$

Try $\hat{\theta} = \bar{X}^3 - 3\bar{X}\sigma^2/N$. This is unbiased and is a function of the sufficient statistic. Thus, it is the MVUE estimator.

$$12) \quad g_1(u) = \frac{1}{N}u \Rightarrow E\left[\frac{1}{N}\sum_n X(n)\right] = A$$

$$g_2(u) = \frac{1}{N}u^{1/3} \Rightarrow E\left[\frac{1}{N}\left(\sum_n X(n)\right)^3\right]^{1/3} = A$$

Also, note that $T_2 = T_1^3$, which is a one-to-one transformation. For T_3 there is no function that would make it unbiased. Also, T_3 is not a one-to-one transformation of T_1 .

$$13) \quad p(\underline{x}; \theta) = e^{-\left(\sum_n X(n) - \theta\right)} u(\min X(n) - \theta)$$

$$= \underbrace{e^{-\sum_n X(n)}}_{h(\underline{x})} \underbrace{e^{N\theta} u(\min X(n) - \theta)}_{g(T(\underline{x}), \theta)}$$

where $T(x) = \min x(n)$ is the sufficient statistic. To find the MLE we proceed as in Example 5.8.

$$\begin{aligned} P_r \{T \leq z\} &= 1 - P_r \{T \geq z\} \\ &= 1 - P_r \{x(0) \geq z, \dots, x(N-1) \geq z\} \\ &= 1 - \prod_{n=0}^{N-1} P_r \{x(n) \geq z\} \\ &= 1 - P_r^N \{x(n) \geq z\} \end{aligned}$$

$$p_T(z) = \frac{d P_r \{T \leq z\}}{dz} = -N P_r \{x(n) \geq z\}^{N-1} \cdot \frac{d P_r \{x(n) \geq z\}}{dz}$$

$$\begin{aligned} \text{But } \frac{d P_r \{x(n) \geq z\}}{dz} &= \frac{d [1 - P_r \{x(n) \leq z\}]}{dz} \\ &= - \frac{d P_r \{x(n) \leq z\}}{dz} = \begin{cases} -e^{-(z-\theta)} & z > \theta \\ 0 & z < \theta \end{cases} \end{aligned}$$

and

$$P_r \{x(n) \geq z\} = \begin{cases} 1 - \int_0^z e^{-(x-\theta)} dx & \text{for } z > \theta \\ 1 & \text{for } z < \theta \end{cases}$$

$$\begin{aligned} \text{For } z > \theta &= 1 + e^0 e^{-x} \Big|_0^z = 1 + e^{-(z-\theta)} - 1 \\ &= e^{-(z-\theta)} \end{aligned}$$

$$\begin{aligned}
 p_T(z) &= \begin{cases} 0 & z < \theta \\ -N [e^{-(z-\theta)^N}]' [-e^{-(z-\theta)^N}] & z > \theta \end{cases} \\
 &= \begin{cases} N e^{-N(z-\theta)} & z > \theta \\ 0 & \text{otherwise} \end{cases}
 \end{aligned}$$

$$\begin{aligned}
 E(T) &= \int_{\theta}^{\infty} z N e^{-N(z-\theta)} dz \\
 &= N e^{N\theta} \int_{\theta}^{\infty} z e^{-Nz} dz \\
 &= N e^{N\theta} \left(-\frac{z}{N} e^{-Nz} - \frac{1}{N^2} e^{-Nz} \right) \Big|_{\theta}^{\infty} \\
 &= N e^{N\theta} \left(\frac{\theta}{N} e^{-N\theta} + \frac{1}{N^2} e^{-N\theta} \right) = \theta + \frac{1}{N}
 \end{aligned}$$

$\Rightarrow \hat{\theta} = T - 1/N = \min x_i - \frac{1}{N}$ is the MVU estimator.

$$\begin{aligned}
 14) \ a) \ p(x; \mu) &= \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}x^2 + x\mu - \frac{1}{2}\mu^2} \\
 &= e^{\underset{\substack{\uparrow \\ A(\mu) B(x)}}{x\mu} - \frac{1}{2}x^2 + \left(-\frac{1}{2}\mu^2 + \ln \frac{1}{\sqrt{2\pi}} \right)} \\
 &\quad \quad \quad \uparrow \quad \quad \quad \uparrow \quad \quad \quad \uparrow \\
 &\quad \quad \quad A(\mu) B(x) \quad C(x) \quad D(\mu)
 \end{aligned}$$

$$\begin{aligned}
 b) \ p(x; \sigma^2) &= \frac{x}{\sigma^2} e^{-\frac{1}{2}x^2/\sigma^2} u(x) \\
 &= e^{\ln x/\sigma^2 - \frac{1}{2}x^2/\sigma^2 + \ln u(x)} \\
 &= e^{\underbrace{-\frac{1}{2}x^2/\sigma^2}_{A(\sigma^2) B(x)} + \underbrace{\ln x u(x)}_{C(x)} - \underbrace{\ln \sigma^2}_{D(\sigma^2)}}
 \end{aligned}$$

$$\begin{aligned}
 c) \quad p(x; \lambda) &= \lambda e^{-\lambda x} u(x) \\
 &= \underbrace{e^{-\lambda x}}_{A(\lambda)B(x)} \underbrace{u(x)}_{C(x)} \underbrace{\lambda}_{D(\lambda)}
 \end{aligned}$$

$$\begin{aligned}
 15) \quad p(\underline{x}; \theta) &= \prod_{n=0}^{N-1} e^{A(\theta)B(x[n]) + C(x[n]) + D(\theta)} \\
 &= e^{A(\theta) \sum_n B(x[n]) + \sum_n C(x[n]) + ND(\theta)} \\
 &= \underbrace{e^{A(\theta) \sum_n B(x[n]) + ND(\theta)}}_{g(T(\underline{x}), \theta)} \underbrace{e^{\sum_n C(x[n])}}_{h(\underline{x})}
 \end{aligned}$$

$$\text{where } T(\underline{x}) = \sum_{n=0}^{N-1} B(x[n])$$

$$a) \quad B(x) = x \Rightarrow T(\underline{x}) = \sum_n x[n]$$

$$b) \quad B(x) = x^2 \Rightarrow T(\underline{x}) = \sum_n x^2[n]$$

$$c) \quad B(x) = x \Rightarrow T(\underline{x}) = \sum_n x[n]$$

Need only make $T(\underline{x})$ unbiased

$$a) \quad \hat{\mu} = 1/N \sum_n x[n]$$

$$\begin{aligned}
 b) \quad E(x^2) &= \int_0^{\infty} \frac{x^2}{\sigma^2} e^{-\frac{1}{2} x^2 / \sigma^2} dx \\
 &= 2\sigma^2
 \end{aligned}$$

$$\Rightarrow \hat{\sigma}^2 = \frac{1}{2N} \sum_n x^2 \ln x$$

$$c) E(x) = \int_0^\infty x \lambda e^{-\lambda x} dx = 1/\lambda$$

It is not obvious how to make T unbiased for this PDF. However, if we reparameterize the PDF by $\theta = 1/\lambda$, the MVL of θ is easily found.

$$1b) p(\underline{x}; \underline{\theta}) = \frac{1}{(2\pi)^{N/2} \det^{1/2}(\underline{c})} e^{-\frac{1}{2}(\underline{x}-\underline{\mu})^T \underline{c}^{-1}(\underline{x}-\underline{\mu})}$$

$$\text{If we have } \underline{c} = \begin{pmatrix} a & \underline{b}^T \\ \underline{b} & \underline{I} \end{pmatrix} \quad \begin{aligned} a &= N-1 + \sigma^2 \\ \underline{b} &= -\underline{1} \end{aligned}$$

$$\begin{aligned} \det(\underline{c}) &= \det(\underline{I}) \det(a - \underline{b}^T \underline{I}^{-1} \underline{b}) \\ &= a - \underline{b}^T \underline{b} = N-1 + \sigma^2 - (-\underline{1})^T(-\underline{1}) \\ &= N-1 + \sigma^2 - (N-1) = \sigma^2 \end{aligned}$$

$$\begin{aligned} \underline{c}^{-1} &= \begin{bmatrix} (a - \underline{b}^T \underline{b})^{-1} & - (a - \underline{b}^T \underline{b})^{-1} \underline{b}^T \\ - \underline{b} (a - \underline{b}^T \underline{b})^{-1} & (\underline{I} - \frac{\underline{b} \underline{b}^T}{a})^{-1} \end{bmatrix} \\ &= \begin{bmatrix} \frac{1}{\sigma^2} & \frac{1}{\sigma^2} \underline{1}^T \\ \frac{1}{\sigma^2} \underline{1} & (\underline{I} - \frac{\underline{1} \underline{1}^T}{N-1 + \sigma^2})^{-1} \end{bmatrix} \end{aligned}$$

$$\text{But } (\underline{I} - \frac{\underline{1} \underline{1}^T}{N-1 + \sigma^2})^{-1} = \underline{I} + \frac{\frac{\underline{1} \underline{1}^T}{N-1 + \sigma^2}}{1 - \frac{\underline{1}^T \underline{1}}{N-1 + \sigma^2}}$$

$$= \underline{I} + \frac{\underline{1}\underline{1}^T}{\sigma^2}$$

$$\underline{C}^{-1} = \frac{1}{\sigma^2} \begin{pmatrix} 1 & \underline{1}^T \\ \underline{1} & \sigma^2 \underline{I} + \underline{1}\underline{1}^T \end{pmatrix}$$

$$(\underline{x} - \underline{\mu})^T \underline{C}^{-1} (\underline{x} - \underline{\mu}) =$$

$$\frac{1}{\sigma^2} \left[\underbrace{(x[0] - N\mu \quad x[1] \dots x[N-1])}_{\underline{x}'^T} \begin{pmatrix} 1 & \underline{1}^T \\ \underline{1} & \sigma^2 \underline{I} + \underline{1}\underline{1}^T \end{pmatrix} \begin{pmatrix} x[0] - N\mu \\ \underline{x}' \end{pmatrix} \right]$$

$$= \frac{1}{\sigma^2} \left\{ (x[0] - N\mu)^2 + (x[0] - N\mu) \underline{1}^T \underline{x}' + (x[0] - N\mu) \underline{x}'^T \underline{1} + \sigma^2 \underline{x}'^T \underline{x}' + (\underline{x}'^T \underline{1})^2 \right\}$$

$$= \frac{1}{\sigma^2} \left\{ (x[0] - N\mu + \underline{1}^T \underline{x}')^2 + \sigma^2 \underline{x}'^T \underline{x}' \right\}$$

$$= \frac{1}{\sigma^2} \left\{ \left(\sum_{n=0}^{N-1} x[n] - N\mu \right)^2 + \sigma^2 \sum_{n=1}^{N-1} x^2[n] \right\}$$

$$= \frac{N^2}{\sigma^2} \left(\bar{x} - \mu \right)^2 + \sum_{n=1}^{N-1} x^2[n]$$

$$p(\underline{x}; \underline{\theta}) = \underbrace{\frac{1}{(2\pi)^{N/2} \sigma}}_{g(T(\underline{x}), \underline{\theta})} e^{-\frac{N^2}{2\sigma^2} (\bar{x} - \mu)^2} \underbrace{e^{-\frac{1}{2} \sum_{n=1}^{N-1} x^2[n]}}_{h(\underline{x})}$$

$\Rightarrow \bar{x}$ is a sufficient statistic for $\underline{\theta}$.

$$17) \quad a) \quad p(\underline{x}; A) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_{n=0}^{N-1} (x(n) - A \cos 2\pi f_0 n)^2}$$

$$\text{But } \sum_n (x(n) - A \cos 2\pi f_0 n)^2 =$$

$$\sum_n x^2(n) - 2A \sum_n x(n) \cos 2\pi f_0 n + A^2 \sum_n \cos^2 2\pi f_0 n$$

$$p(\underline{x}; A) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} (A^2 \sum_n \cos^2 2\pi f_0 n - 2A \sum_n x(n) \cos 2\pi f_0 n)}$$

$$= \underbrace{e^{-\frac{1}{2\sigma^2} \sum_n x^2(n)}}_{h(\underline{x})} \underbrace{e^{-\frac{1}{2\sigma^2} (A^2 \sum_n \cos^2 2\pi f_0 n - 2A \sum_n x(n) \cos 2\pi f_0 n)}}_{g(T(\underline{x}), A)}$$

where $T(\underline{x}) = \sum_{n=0}^{N-1} x(n) \cos 2\pi f_0 n$ is the sufficient statistic

$$E(T(\underline{x})) = \sum_{n=0}^{N-1} A \cos^2 2\pi f_0 n$$

$$\Rightarrow \hat{A} = \frac{\sum_{n=0}^{N-1} x(n) \cos 2\pi f_0 n}{\sum_{n=0}^{N-1} \cos^2 2\pi f_0 n}$$

$$b) \quad \text{Let } \underline{\theta} = \begin{bmatrix} A \\ \sigma^2 \end{bmatrix}$$

From part a we have

$$T(\underline{x}) = \begin{bmatrix} \sum_n x(n) \cos 2\pi f_0 n \\ \sum_n x^2(n) \end{bmatrix}$$

is a sufficient statistic

To make $T_1(x)$ unbiased let

$$\hat{A} = \frac{\sum_n x[n] \cos 2\pi f_0 n}{\sum_n \cos^2 2\pi f_0 n} \quad \text{as before.}$$

To make $T_2(x)$ unbiased:

$$\begin{aligned} E(T_2(x)) &= \sum_n E(x^2[n]) \\ &= \sum_n E[(A \cos 2\pi f_0 n + w[n])^2] \end{aligned}$$

$$= \sum_n A^2 \cos^2 2\pi f_0 n + N \sigma^2$$

From Example 5.11 we expect that we will have to subtract out the squared mean to generate an unbiased estimator of σ^2 .

But

$$\begin{aligned} E(\hat{A}^2) &= \frac{E\left[\left(\sum_n x[n] \cos 2\pi f_0 n\right)^2\right]}{\left(\sum_n \cos^2 2\pi f_0 n\right)^2} \\ &= \frac{\sum_{m,n} E(x[m] x[n]) \cos 2\pi f_0 m \cos 2\pi f_0 n}{\left(\sum_n \cos^2 2\pi f_0 n\right)^2} \\ &= \frac{\sum_{m,n} \left(A^2 \cos 2\pi f_0 m \cos 2\pi f_0 n + \sigma^2 \delta_{mn} \right) \cdot \begin{matrix} (\cos 2\pi f_0 m) \\ (\cos 2\pi f_0 n) \end{matrix}}{\left(\sum_n \cos^2 2\pi f_0 n\right)^2} \end{aligned}$$

$$= \frac{A^2 \left(\sum_n \cos^2 2\pi f_0 n \right)^2 + \sigma^2 \sum_n \cos^2 2\pi f_0 n}{\left(\sum_n \cos^2 2\pi f_0 n \right)^2}$$

$$= A^2 + \frac{\sigma^2}{\sum_n \cos^2 2\pi f_0 n}$$

so that if $T_2'(x) = T_2(x) - \sum_n \cos^2 2\pi f_0 n \hat{A}^2$

$$\begin{aligned} E(T_2'(x)) &= A^2 \sum_n \cos^2 2\pi f_0 n + N \sigma^2 \\ &\quad - A^2 \sum_n \cos^2 2\pi f_0 n \cdot \sigma^2 \\ &= (N-1) \sigma^2 \end{aligned}$$

$\Rightarrow$ Let $T_2''(x) = \frac{1}{N-1} T_2'(x)$

$$= \frac{1}{N-1} \left[\sum_n x^2(n) - \sum_n \cos^2 2\pi f_0 n \hat{A}^2 \right]$$

$\therefore \hat{\theta} = \begin{bmatrix} \hat{A} \\ \sigma^2 \end{bmatrix}$

$$= \begin{bmatrix} \frac{\sum_{n=0}^{N-1} x(n) \cos 2\pi f_0 n}{\sum_{n=0}^{N-1} \cos^2 2\pi f_0 n} \\ \frac{1}{N-1} \left[\sum_{n=0}^{N-1} x^2(n) - \hat{A}^2 \sum_{n=0}^{N-1} \cos^2 2\pi f_0 n \right] \end{bmatrix}$$

$$18) \quad p(x(n)) = \begin{cases} \frac{1}{\theta_2 - \theta_1} & \theta_1 \leq x(n) \leq \theta_2 \\ 0 & \text{otherwise} \end{cases}$$

$$p(\underline{x}; \underline{\theta}) = \begin{cases} \frac{1}{(\theta_2 - \theta_1)^N} & \text{all } x(n) \text{ satisfy} \\ & \theta_1 \leq x(n) \leq \theta_2 \\ 0 & \text{otherwise} \end{cases}$$

Alternatively, for the PDF to be nonzero

$\min x(n) \geq \theta_1, \max x(n) \leq \theta_2$ so that

$$p(\underline{x}; \underline{\theta}) = \underbrace{\frac{1}{(\theta_2 - \theta_1)^N} u(\min x(n) - \theta_1) u(\max x(n) - \theta_2)}_{g(T(\underline{x}), \underline{\theta})}$$

$$\underbrace{\quad}_{h(\underline{x})}$$

$$\Rightarrow T(\underline{x}) = \begin{bmatrix} \min x(n) \\ \max x(n) \end{bmatrix} \text{ is a sufficient statistic}$$

$$\begin{aligned} 19) \quad & (\underline{X} - \underline{H}\hat{\underline{\theta}})^T (\underline{X} - \underline{H}\hat{\underline{\theta}}) + (\hat{\underline{\theta}} - \underline{\theta})^T \underline{H}^T \underline{H} (\hat{\underline{\theta}} - \underline{\theta}) \\ &= (\underline{X} - \underline{H}(\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{X})^T (\underline{X} - \underline{H}(\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{X}) \\ &\quad + (\hat{\underline{\theta}} - (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{X})^T \underline{H}^T \underline{H} (\hat{\underline{\theta}} - (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{X}) \\ &= \underline{X}^T (\underline{I} - \underline{H}(\underline{H}^T \underline{H})^{-1} \underline{H}^T) (\underline{I} - \underline{H}(\underline{H}^T \underline{H})^{-1} \underline{H}^T) \underline{X} \\ &\quad + \hat{\underline{\theta}}^T \underline{H}^T \underline{H} \hat{\underline{\theta}} - \hat{\underline{\theta}}^T \underline{H}^T \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{X} \\ &\quad - \underline{X}^T \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{H} \hat{\underline{\theta}} + \underline{X}^T \underline{H}^T (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{X} \\ &= \underline{X}^T \underline{X} - \underline{X}^T \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{X} + \hat{\underline{\theta}}^T \underline{H}^T \underline{H} \hat{\underline{\theta}} - 2 \hat{\underline{\theta}}^T \underline{H}^T \underline{X} \\ &\quad + \underline{X}^T \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{X} = (\underline{X} - \underline{H}\hat{\underline{\theta}})^T (\underline{X} - \underline{H}\hat{\underline{\theta}}) \end{aligned}$$

$$\begin{aligned}
 p(\underline{x}; \underline{\theta}) &= \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} (\underline{x} - H\underline{\theta})^T (\underline{x} - H\underline{\theta})} \\
 &= \underbrace{\frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} (\underline{\theta} - \hat{\underline{\theta}})^T H^T H (\underline{\theta} - \hat{\underline{\theta}})}}_{g(T(\underline{x}), \underline{\theta})} \underbrace{e^{-\frac{1}{2\sigma^2} (\underline{x} - H\hat{\underline{\theta}})^T (\underline{x} - H\hat{\underline{\theta}})}}_{h(\underline{x})}
 \end{aligned}$$

where $T(\underline{x}) = \hat{\underline{\theta}}$ = sufficient statistic
 Since we already know that $\hat{\underline{\theta}}$ is unbiased,
 it is the MVU estimator. This is not
 unexpected since we saw in Chapter 3
 that $\hat{\underline{\theta}}$ is efficient.

Chapter 6

$$1) \quad \hat{A} = (\underline{H}^T \underline{C}^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}^{-1} \underline{x}$$

$$\text{where } \underline{H} = \begin{bmatrix} 1 \\ r \\ \vdots \\ r^{N-1} \end{bmatrix} \quad \underline{C} = \sigma^2 \underline{I}$$

$$\begin{aligned} \hat{A} &= \left(\frac{1}{\sigma^2} \sum_{n=0}^{N-1} r^{2n} \right)^{-1} \frac{1}{\sigma^2} \sum_{n=0}^{N-1} x[n] r^n \\ &= \frac{\sum_{n=0}^{N-1} x[n] r^n}{\sum_{n=0}^{N-1} r^{2n}} \end{aligned}$$

$$\text{var}(\hat{A}) = \frac{1}{\underline{H}^T \underline{C}^{-1} \underline{H}} = \frac{\sigma^2}{\sum_{n=0}^{N-1} r^{2n}}$$

$$\text{var}(\hat{A}) \rightarrow 0 \quad \text{if } |r| \geq 1.$$

$$2) \quad \text{var}(\hat{A}) = \frac{1}{\sum_{n=0}^{N-1} 1/\sigma_n^2}$$

$$\text{If } \sigma_n^2 = n+1, \quad \text{var}(\hat{A}) = \frac{1}{\sum_{n=0}^{N-1} 1/(n+1)}$$

As $N \rightarrow \infty$, $\sum_{n=0}^{N-1} \frac{1}{n+1} \rightarrow \infty$ since this is a harmonic series $\Rightarrow \text{var}(\hat{A}) \rightarrow 0$

$$\text{If } \sigma_n^2 = (n+1)^2, \quad \text{as } N \rightarrow \infty \quad \sum_{n=0}^{N-1} \frac{1}{(n+1)^2} \rightarrow \text{Constant} \\ \Rightarrow \text{var}(\hat{A}) \not\rightarrow 0.$$

In this case the noise samples have such a large variance that the estimator

variance does not go to zero.

$$3) \quad \hat{A} = \frac{\underline{1}^T \underline{C}^{-1} \underline{x}}{\underline{1}^T \underline{C}^{-1} \underline{1}}$$

$$\underline{C}^{-1} = \frac{1}{\sigma^2} \begin{bmatrix} \underline{B} & & 0 \\ & \underline{B} & \\ 0 & \ddots & \underline{B} \end{bmatrix} \quad \text{where } \underline{B} = \frac{\begin{bmatrix} 1 & -\rho \\ -\rho & 1 \end{bmatrix}}{1-\rho^2}$$

$$\underline{1}^T \underline{C}^{-1} \underline{1} = \text{Sum of all elements in } \underline{C}^{-1}$$

$$= \frac{N/2}{\sigma^2} \frac{2-2\rho}{1-\rho^2} = \frac{N}{\sigma^2(1+\rho)}$$

$$\text{Let } \underline{x} = [\underline{x}_1^T \underline{x}_2^T \dots \underline{x}_{N/2}^T]^T \quad \text{where each } \underline{x}_i \text{ is } 2 \times 1.$$

$$\underline{1}^T \underline{C}^{-1} \underline{x} = \frac{1}{\sigma^2} \underline{1}^T \begin{bmatrix} \underline{B} \underline{x}_1 \\ \vdots \\ \underline{B} \underline{x}_{N/2} \end{bmatrix}$$

$$= \frac{1}{\sigma^2} \sum_{i=1}^{N/2} \underset{\substack{\uparrow \\ 2 \times 1}}{\underline{1}^T \underline{B} \underline{x}_i}$$

$$\text{But } \underline{B}^T \underline{1} = \frac{\begin{bmatrix} 1-\rho \\ 1-\rho \end{bmatrix}}{1-\rho^2} \Rightarrow \underline{1}^T \underline{B} \underline{x}_i = \underline{x}_i^T \underline{B}^T \underline{1} = \frac{(1-\rho)[\underline{x}_i]_1 + (1-\rho)[\underline{x}_i]_2}{1-\rho^2}$$

$$= \frac{[\underline{x}_i]_1 + [\underline{x}_i]_2}{1+\rho}$$

$$\underline{1}^T \underline{C}^{-1} \underline{x} = \frac{1}{\sigma^2} \frac{1}{1+\rho} \sum_{i=1}^{N/2} ([\underline{x}_i]_1 + [\underline{x}_i]_2)$$

$$= \frac{1}{\sigma^2} \frac{1}{1+\rho} \sum_{n=0}^{N-1} x(n)$$

$$\hat{A} = \frac{\frac{1}{\sigma^2} \frac{N \bar{x}}{1+p}}{\frac{N}{\sigma^2(1+p)}} = \bar{x}$$

$$\text{var}(\hat{A}) = \frac{1}{\frac{N}{\sigma^2(1+p)}} = \frac{\sigma^2(1+p)}{N}$$

Since the subvectors $\underline{x}_1, \underline{x}_2, \dots, \underline{x}_{N/2}$ are uncorrelated, we average them. For a single subvector we also average the samples (see Probs. 3.9 and 4.11). Hence, we obtain $\bar{x}$. The variance is σ^2/N for $p=0$ (our usual case), $2\sigma^2/N$ for $p \rightarrow 1$ since the samples of each subvector are equal and hence we have only $N/2$ uncorrelated samples, and $\rightarrow 0$ for $p \rightarrow -1$ since then the noise samples cancel (see Probs. 3.9 and 4.11).

4) In either case we have the model

$$\underline{x} = \underline{1}u + \underline{w} \quad \text{where } E(\underline{w}) = \underline{0}$$

$$\text{and } E(\underline{w}\underline{w}^T) = \text{var}(\underline{w}) \underline{I}$$

The BLUE for each case is

$$\hat{u} = \frac{\underline{1}^T \underline{C}^{-1} \underline{x}}{\underline{1}^T \underline{C}^{-1} \underline{1}} = \frac{\underline{1}^T \underline{x}}{\underline{1}^T \underline{1}} = \bar{x}$$

But in the Gaussian case the BLUE is also the MVU estimator - not so for the Laplacian PDF.

$$(5) \quad E(x) = \int_0^{\infty} x \frac{1}{\sqrt{2\pi} x} e^{-\frac{1}{2}(\ln x - \theta)^2} dx$$

$$\text{Let } y = \ln x \quad dy/dx = 1/x \Rightarrow dx = e^y dy$$

$$\begin{aligned} E(x) &= \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(y-\theta)^2} e^y dy \\ &= \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(y^2 - 2y(\theta+1) + (\theta+1)^2)} \\ &\quad \cdot e^{-\frac{1}{2}(\theta^2 - (\theta+1)^2)} dy \\ &= \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(y - (\theta+1))^2} dy \\ &\quad \cdot e^{-\frac{1}{2}(-2\theta-1)} \\ &= e^{\theta+1/2} \end{aligned}$$

$$\text{Now let } y = \ln x, \quad dy/dx = 1/x = e^{-y}$$

$$\begin{aligned} p(y) &= \frac{p(x(y))}{|dy/dx|} = \frac{\frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(y-\theta)^2}}{e^{-y}} \\ &= \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(y-\theta)^2} \end{aligned}$$

$$\Rightarrow y \sim N(\theta, 1)$$

The BLUE is $\hat{\theta} = \frac{1}{N} \sum_{n=0}^{N-1} y[n]$

$$= \frac{1}{N} \sum_{n=0}^{N-1} \ln x[n]$$

b) For $\hat{\theta}$ to be unbiased

$$\begin{aligned} E(\hat{\theta}) &= E\left(\sum_n a_n x[n] + b\right) \\ &= \sum_n a_n (0.5) + \beta + b \\ &= \theta \end{aligned}$$

$$\Rightarrow \sum_n a_n 0.5 = 1, \quad \beta \sum_n a_n = -b$$

$$\begin{aligned} \text{or } \hat{\theta} &= \sum_n a_n x[n] - \beta \sum_n a_n \\ &= \sum_n a_n (x[n] - \beta) \end{aligned}$$

Let $x'[n] = x[n] - \beta$. Then, we have the same problem as before. Thus,

$$\hat{\theta} = \frac{\underline{y}^T \underline{C}^{-1} \underline{x}'}{\underline{y}^T \underline{C}^{-1} \underline{y}} = \frac{\underline{y}^T \underline{C}^{-1} (\underline{x} - \beta \underline{1})}{\underline{y}^T \underline{C}^{-1} \underline{y}}$$

and since the covariance for $\underline{x} - \beta \underline{1}$ is $\underline{C}$,

$$\text{var}(\hat{\theta}) = \frac{1}{\underline{y}^T \underline{C}^{-1} \underline{y}}$$

7) $\hat{A} = \frac{\underline{y}^T \underline{C}^{-1} \underline{x}}{\underline{y}^T \underline{C}^{-1} \underline{y}}$

Assume $\underline{C} \underline{y} = \lambda \underline{y}$. Then, $\underline{C}^{-1} \underline{y} = 1/\lambda \underline{y}$

$$\hat{A} = \frac{\underline{x}^T \underline{C}^{-1} \underline{y}}{\underline{y}^T \underline{C}^{-1} \underline{y}} = \frac{\underline{x}^T \underline{1}/\lambda \underline{y}}{\underline{y}^T \underline{1}/\lambda \underline{y}} = \frac{\underline{x}^T \underline{y}}{\underline{y}^T \underline{y}}$$

$$\text{var}(\hat{A}) = \frac{1}{\underline{y}^T \underline{C}^{-1} \underline{y}} = \frac{\lambda}{\underline{y}^T \underline{y}}$$

In this case we obtain same result as if $\underline{C} = \sigma^2 \underline{I}$. We do not need a prewhitener.

$$\begin{aligned} 8) \quad \text{var}(\hat{A}) &= \frac{1}{\underline{y}^T \underline{C}^{-1} \underline{y}} \\ &= \frac{1}{\left(\sum_i \alpha_i \underline{v}_i \right)^T \underline{C}^{-1} \left(\sum_j \alpha_j \underline{v}_j \right)} \\ &= \frac{1}{\sum_i \sum_j \alpha_i \alpha_j \underbrace{\underline{v}_i^T \underline{C}^{-1} \underline{v}_j}_{\underline{v}_i^T \underline{1}/\lambda_j \underline{v}_j}} \\ &= \frac{1}{\sum_i \alpha_i^2 / \lambda_i} \end{aligned}$$

$$\begin{aligned} \Sigma = \underline{y}^T \underline{y} &= \left(\sum_i \alpha_i \underline{v}_i \right)^T \left(\sum_j \alpha_j \underline{v}_j \right) \\ &= \sum_i \sum_j \alpha_i \alpha_j \underbrace{\underline{v}_i^T \underline{v}_j}_{\delta_{ij}} = \sum_i \alpha_i^2 \end{aligned}$$

Must minimize $\frac{1}{\sum_i \alpha_i^2 / \lambda_i}$ subject
to constraint $\sum_i \alpha_i^2 = \epsilon_0$. Equivalently,
we must maximize $\sum_i \alpha_i^2 / \lambda_i$. Using
Lagrangian multipliers

$$F = \sum_i \alpha_i^2 / \lambda_i + \lambda \left(\sum_i \alpha_i^2 - \epsilon_0 \right)$$

$$\frac{\partial F}{\partial \alpha_k} = \frac{2\alpha_k}{\lambda_k} + \lambda 2\alpha_k = 0$$

$$\Rightarrow \alpha_k = 0 \text{ or}$$

$$\lambda = -1/\lambda_k \text{ all } k$$

Clearly, we cannot have $\alpha_k = 0$ for all k , since
then constraint could not be satisfied.
Since the eigenvalues are distinct, we
also cannot have $\lambda = -1/\lambda_k$ for all k .
Thus, we must have

$$\alpha_k = 0 \text{ except for } k=j$$

$$\lambda = -1/\lambda_j \text{ and } \alpha_j \neq 0.$$

Hence, $\underline{s} = \alpha_j \underline{v}_j$. To determine which
eigenvector to use

$$\text{var}(\hat{A}) = \frac{1}{\alpha_j^2 / \lambda_j} = \lambda_j / \alpha_j^2$$

And since $\epsilon_0 = \alpha_j^2$, $\text{var}(\hat{A}) = \lambda_j / \epsilon_0$ so that λ_j should be the minimum eigenvalue. Hence, the optimal signal is

$$\underline{s} = c \underline{v}_{\text{MIN}} = \sqrt{\epsilon_0} \underline{v}_{\text{MIN}}$$

where $\underline{v}_{\text{MIN}}$ is the eigenvector associated with the smallest eigenvalue. Intuitively, we place signal along direction where there is the least amount of noise.

9) $\theta = A$

$$\underline{s} = [1 \cos 2\pi f_1 \dots \cos 2\pi f_1 (N-1)]^T$$

$$\hat{A} = \frac{\underline{s}^T \underline{C}^{-1} \underline{x}}{\underline{s}^T \underline{C}^{-1} \underline{s}} = \frac{\underline{s}^T \underline{x}}{\underline{s}^T \underline{s}}$$

$$= \frac{\sum_{n=0}^{N-1} x[n] \cos 2\pi f_1 n}{\sum_{n=0}^{N-1} \cos^2 2\pi f_1 n}$$

This is a scaled Fourier coefficient since

$$\hat{A} = \frac{N/2}{\sum_n \cos^2 2\pi f_1 n} \underbrace{\frac{2}{N} \sum_n x[n] \cos 2\pi f_1 n}_{\text{Fourier coefficient}}$$

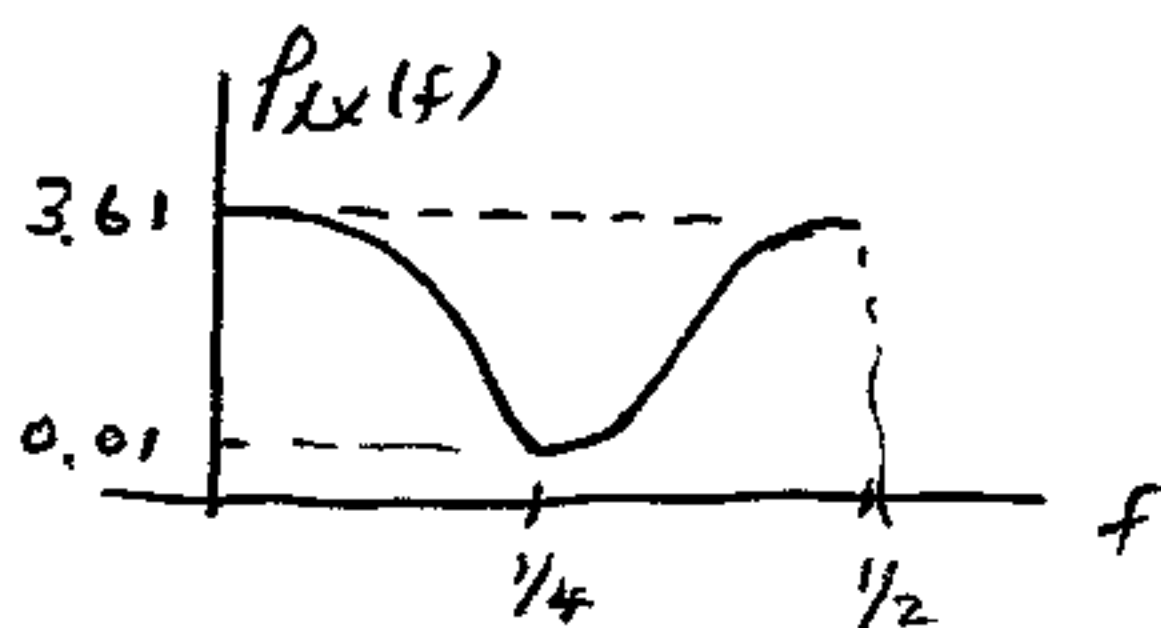
and for large N and f_1 not near 0 or $1/2$ the scale factor is one.

$$\begin{aligned} \text{var}(\hat{A}) &= \frac{1}{\underline{s}^T \underline{C}^{-1} \underline{s}} = \frac{\sigma^2}{\sum_{n=0}^{N-1} s^2(n)} \\ &= \frac{\sigma^2}{\sum_{n=0}^{N-1} \cos^2 2\pi f_1 n} \geq \frac{\sigma^2}{N} \end{aligned}$$

Since $\sum_{n=0}^{N-1} \cos^2 2\pi f_1 n \leq N$

$\Rightarrow$ variance is minimized if $f_1 = 0$.
Note that this choice maximizes the signal energy.

$$\begin{aligned} 10) \quad P_{xx}(f) &= r_{xx}(0) + r_{xx}(2) e^{-j4\pi f} + r_{xx}(2) e^{j4\pi f} \\ &= 1.81 + 2(0.9) \cos 4\pi f \\ &= 1.81 + 1.8 \cos 4\pi f \end{aligned}$$

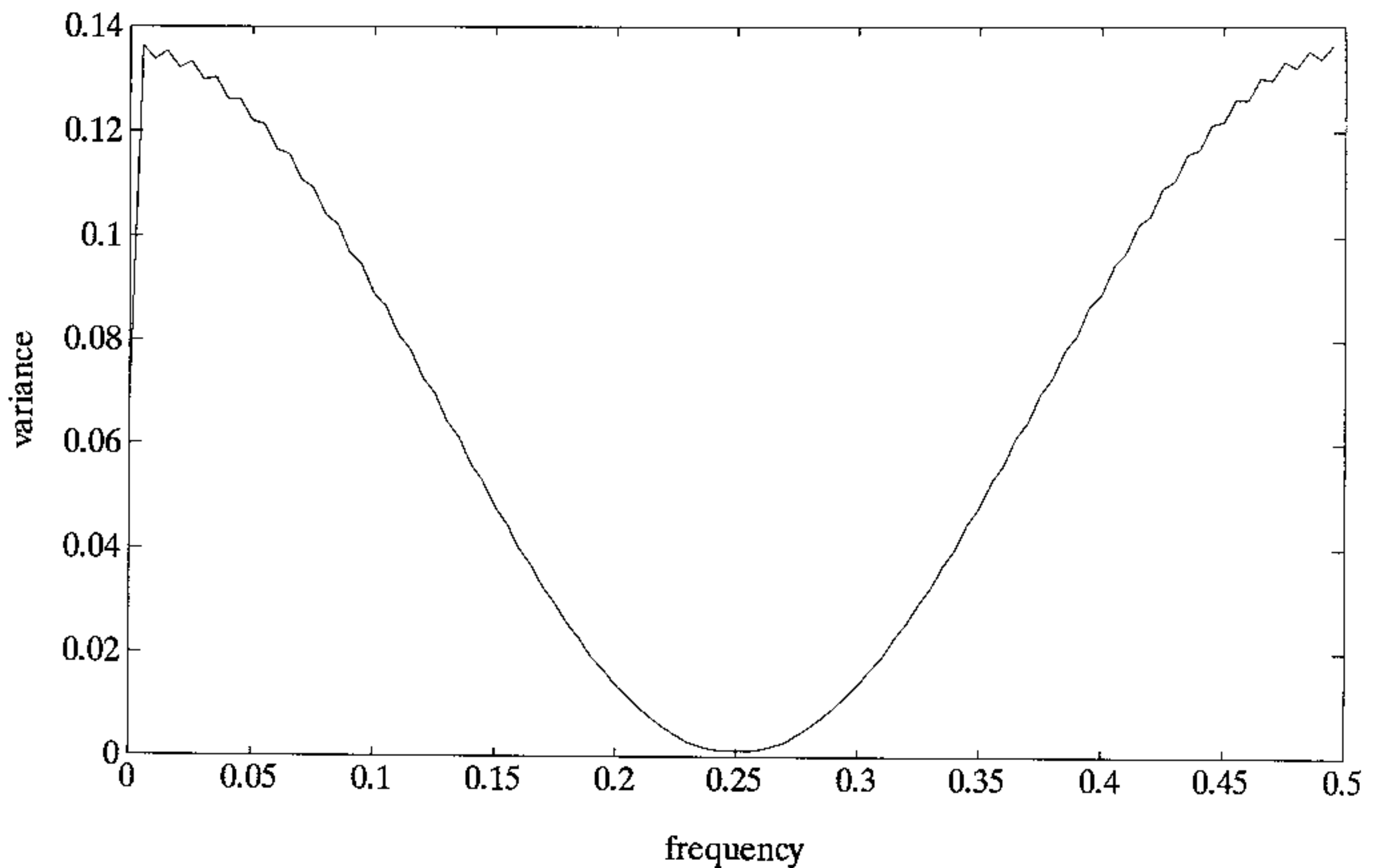


Need to minimize $\frac{1}{\underline{s}^T \underline{C}^{-1} \underline{s}} = \text{var}(\hat{A})$
where $\underline{s} = [1 \cos 2\pi f_1 \dots \cos 2\pi f_1 (N-1)]^T$
and

$$\underline{C} = \begin{bmatrix} r_{xx}(0) & r_{xx}(1) & \dots & r_{xx}(N-1) \\ r_{xx}(1) & r_{xx}(0) & & \\ & & \ddots & \\ & & & r_{xx}(0) \end{bmatrix}$$

$$= \begin{bmatrix} 1.81 & 0 & 0.9 & 0 & \dots & 0 \\ 0 & 1.81 & 0 & 0.9 & \dots & 0 \\ & & \ddots & & & \\ & & & & & 1.81 \end{bmatrix}$$

As seen in the following graph, the variance is minimized for $f_1 = 0.25$ or where the PSD is minimum.



$$11) \quad H(e^{j2\pi f}) = \sum_{n=0}^{N-1} h[n] e^{-j2\pi f n}$$

$$H(e^{j0}) = \sum_{n=0}^{N-1} h[n] = 1$$

The noise power at the output at $n = N-1$ is

$$\begin{aligned} & E \left[\left(\sum_k h[k] w[N-1-k] \right)^2 \right] \\ &= \sum_k \sum_l h[k] h[l] \underbrace{E[w[N-1-k] w[N-1-l]]}_{r_{ww}[k-l]} \\ &= \underline{h}^T \underline{C} \underline{h} \quad \text{where } (C)_{kl} = r_{ww}[k-l] \end{aligned}$$

Thus, to find the FIR filter coefficients we need to minimize $\underline{h}^T \underline{C} \underline{h}$ subject to the constraint $\underline{h}^T \underline{1} = 1$. This is just the BLUE setup so that

$$\underline{h}_{opt} = \frac{\underline{C}^{-1} \underline{1}}{\underline{1}^T \underline{C}^{-1} \underline{1}}$$

The noise power at the output is

$$\underline{h}_{opt}^T \underline{C} \underline{h}_{opt} = \frac{\underline{1}^T \underline{C}^{-1} \underline{C} \underline{C}^{-1} \underline{1}}{(\underline{1}^T \underline{C}^{-1} \underline{1})^2}$$

$$= \frac{1}{\mathbf{1}^T \mathbf{C}^{-1} \mathbf{1}}$$

which is just the variance of $\hat{A}$ since

$$\begin{aligned} \text{var}(\hat{A}) &= E[(\hat{A} - E(\hat{A}))^2] \\ &= E\left[\left(\sum_k h[k] x[N-1-k] - \sum_k h[k] A\right)^2\right] \end{aligned}$$

Since $\sum_k h[k] = 1$

$$\begin{aligned} &= E\left[\left(\sum_k h[k] (x[N-1-k] - A)\right)^2\right] \\ &= E\left[\left(\sum_k h[k] w[N-1-k]\right)^2\right] \end{aligned}$$

The BLUE may be viewed as the output of a linear filter constrained to pass the DC signal and whose coefficients are chosen to minimize the noise at the filter output.

12) Since $\underline{x} = \underline{H}\underline{\theta} + \underline{w}$ and $\underline{B}$ is invertible,

$$\underline{0} = \underline{B}^{-1}(\underline{x} - \underline{b}) \Rightarrow \underline{x} = \underline{H}\underline{B}^{-1}(\underline{x} - \underline{b}) + \underline{w}$$

$$\text{or } \underline{x} = \underline{H}\underline{B}^{-1}\underline{x} - \underline{H}\underline{B}^{-1}\underline{b} + \underline{w}$$

$$\underbrace{\underline{x} + \underline{H}\underline{B}^{-1}\underline{b}}_{\underline{x}'} = \underbrace{\underline{H}\underline{B}^{-1}\underline{x}}_{\underline{u}'} + \underline{w}$$

$$\begin{aligned}
\Rightarrow \hat{\underline{\alpha}} &= (\underline{H}'^T \underline{C}^{-1} \underline{H}')^{-1} \underline{H}'^T \underline{C}^{-1} \underline{x}' \\
&= (\underline{B}^{-1T} \underline{H}^T \underline{C}^{-1} \underline{H} \underline{B}^{-1})^{-1} \underline{B}^{-1T} \underline{H}^T \underline{C}^{-1} (\underline{x} + \underline{H} \underline{B}^{-1} \underline{b}) \\
&= \underline{B} (\underline{H}^T \underline{C}^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}^{-1} (\underline{x} + \underline{H} \underline{B}^{-1} \underline{b}) \\
&= \underline{B} \hat{\underline{\theta}} + \underline{B} \underline{B}^{-1} \underline{b} = \underline{B} \hat{\underline{\theta}} + \underline{b}
\end{aligned}$$

$$\begin{aligned}
13) \quad J &= \underline{x}^T \underline{C}^{-1} \underline{x} - \underline{x}^T \underline{C}^{-1} \underline{H} \underline{\theta} - \underline{\theta}^T \underline{H}^T \underline{C}^{-1} \underline{x} \\
&\quad + \underline{\theta}^T \underline{H}^T \underline{C}^{-1} \underline{H} \underline{\theta} \\
&= \underline{x}^T \underline{C}^{-1} \underline{x} - 2 \underline{\theta}^T \underline{H}^T \underline{C}^{-1} \underline{x} + \underline{\theta}^T \underline{H}^T \underline{C}^{-1} \underline{H} \underline{\theta}
\end{aligned}$$

Using (4.3) we have

$$\frac{\partial J}{\partial \underline{\theta}} = -2 \underline{H}^T \underline{C}^{-1} \underline{x} + 2 \underline{H}^T \underline{C}^{-1} \underline{H} \underline{\theta} = \underline{0}$$

$$\Rightarrow \hat{\underline{\theta}} = (\underline{H}^T \underline{C}^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}^{-1} \underline{x}$$

14) Since $p(W(n))$ is even, $E(W(n)) = 0$.

$$\text{var}(W(n)) = E(W^2(n))$$

$$\begin{aligned}
&= \int_{-\infty}^{\infty} W^2 \left(\frac{1-E}{\sqrt{2\pi\sigma_0^2}} e^{-\frac{1}{2} W^2/\sigma_0^2} \right. \\
&\quad \left. + \frac{E}{\sqrt{2\pi\sigma_1^2}} e^{-\frac{1}{2} W^2/\sigma_1^2} \right) dW
\end{aligned}$$

$$\begin{aligned}
&= (1-\epsilon) \int_{-\infty}^{\infty} W^2 \frac{1}{\sqrt{2\pi\sigma_B^2}} e^{-\frac{1}{2}W^2/\sigma_B^2} dW \\
&\quad + \epsilon \int_{-\infty}^{\infty} W^2 \frac{1}{\sqrt{2\pi\sigma_I^2}} e^{-\frac{1}{2}W^2/\sigma_I^2} dW \\
&= (1-\epsilon)\sigma_B^2 + \epsilon\sigma_I^2
\end{aligned}$$

The BLUE of σ^2 is $\hat{\sigma}^2 \equiv \frac{1}{N} \sum_{n=0}^{N-1} W^2(n)$ (see Section 6.3). This is because the $W(n)$'s are independent and thus $y(n) = W^2(n)$'s are independent. The mean is $E(y(n)) = \sigma^2$ and covariance matrix is $\underline{C} = \text{var}(y(n)) \underline{I}$ so that

$$\hat{\sigma}^2 = \frac{\underline{y}^T \underline{C}^{-1} \underline{y}}{\underline{y}^T \underline{C}^{-1} \underline{y}} = \frac{\underline{1}^T \underline{y}}{\underline{1}^T \underline{1}} = \frac{1}{N} \sum_{n=0}^{N-1} y(n)$$

Now using results from Prob. 6.12

$$\sigma_I^2 = \frac{\sigma^2 - (1-\epsilon)\sigma_B^2}{\epsilon}$$

$$\begin{aligned}
\Rightarrow \hat{\sigma}_I^2 &= \frac{\hat{\sigma}^2 - (1-\epsilon)\sigma_B^2}{\epsilon} \\
&= \frac{\frac{1}{N} \sum_{n=0}^{N-1} W^2(n) - (1-\epsilon)\sigma_B^2}{\epsilon}
\end{aligned}$$

$$15) \quad \underbrace{\underline{X} - \underline{\varepsilon}}_{\underline{X}'} = \underline{H} \underline{\theta} + \underline{W}$$

$$\Rightarrow \hat{\underline{\theta}} = (\underline{H}^T \underline{C}^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}^{-1} \underline{X}'$$

$$= (\underline{H}^T \underline{C}^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}^{-1} (\underline{X} - \underline{\varepsilon})$$

$$16) \quad \hat{A} = (\underline{1}^T \hat{\underline{C}}^{-1} \underline{1})^{-1} \underline{1}^T \hat{\underline{C}}^{-1} \underline{X}$$

$$E(\hat{A}) = (\underline{1}^T \hat{\underline{C}}^{-1} \underline{1})^{-1} \underline{1}^T \hat{\underline{C}}^{-1} \underbrace{E(\underline{X})}_{\underline{A}} = A$$

$\Rightarrow$ unbiased for any $\hat{\underline{C}}$

$$\text{var}(\hat{A}) = E[(\hat{A} - A)^2] = E[(\underline{1}^T \hat{\underline{C}}^{-1} \underline{1})^{-1} \underline{1}^T \hat{\underline{C}}^{-1} \underline{W}]^2]$$

$$= \frac{\underline{1}^T \hat{\underline{C}}^{-1} E(\underline{W} \underline{W}^T) \hat{\underline{C}}^{-1} \underline{1}}{(\underline{1}^T \hat{\underline{C}}^{-1} \underline{1})^2}$$

$$= \frac{\underline{1}^T \hat{\underline{C}}^{-1} \underline{C} \hat{\underline{C}}^{-1} \underline{1}}{(\underline{1}^T \hat{\underline{C}}^{-1} \underline{1})^2}$$

$$\text{If } \underline{C} = \underline{I}, \quad \hat{\underline{C}} = \begin{pmatrix} 1 & 0 \\ 0 & \alpha \end{pmatrix} \Rightarrow \hat{\underline{C}}^{-1} = \begin{pmatrix} 1 & 0 \\ 0 & 1/\alpha \end{pmatrix}$$

$$\text{var}(\hat{A}) = \frac{\underline{1}^T \begin{pmatrix} 1 & 0 \\ 0 & 1/\alpha \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & 1/\alpha \end{pmatrix} \underline{1}}{(\underline{1}^T \begin{pmatrix} 1 & 0 \\ 0 & 1/\alpha \end{pmatrix} \underline{1})^2}$$

$$= \frac{1 + 1/\alpha^2}{(1 + 1/\alpha)^2}$$

Clearly, by the minimum variance property of the BLUE

$$\text{var}(\hat{A})_{\text{MIN}} = 1/2 \quad \text{for } \alpha = 1$$

$$\text{Let } J = \frac{\text{var}(\hat{A})}{\text{var}(\hat{A})_{\text{MIN}}} = \frac{2(1 + 1/\alpha^2)}{(1 + 1/\alpha)^2}$$

$$= \frac{2(1 + \alpha^2)}{(1 + \alpha)^2}$$



For $\alpha \rightarrow 0$ or $\alpha \rightarrow \infty$ we discard one data sample and thus the variance doubles. For $\alpha = 1$ we have the BLUE and hence the minimum variance.

Chapter 7

$$1) \quad p(\underline{x}; A) = \frac{1}{(2\pi A)^{N/2}} e^{-\frac{1}{2A} \sum_n (x[n] - A)^2}$$

$$\frac{\partial \log p}{\partial A} = -\frac{N}{2A} + \frac{1}{A} \sum_n (x[n] - A) + \frac{1}{2A^2} \sum_n (x[n] - A)^2$$

$$\begin{aligned} \frac{\partial^2 \log p}{\partial A^2} &= \frac{N}{2A^2} - \frac{1}{A^2} \sum_n x[n] - \frac{1}{A^2} \sum_n (x[n] - A) \\ &\quad - \frac{1}{A^3} \sum_n (x[n] - A)^2 \end{aligned}$$

$$\begin{aligned} E \left[\frac{\partial^2 \log p}{\partial A^2} \right] &= \frac{N}{2A^2} - \frac{NA}{A^2} - 0 - \frac{1}{A^3} (NA) \\ &= -\frac{N}{2A^2} - \frac{N}{A} \end{aligned}$$

$$I(A) = \frac{N}{2A^2} + \frac{N}{A} \Rightarrow \text{var}(\hat{A}) \geq \frac{1}{\frac{N}{2A^2} + \frac{N}{A}}$$

$$\text{or } \text{var}(\hat{A}) \geq \frac{A^2}{N(A + \frac{1}{2})}$$

$$2) \quad \text{var}(\bar{x}) = \sigma^2/N = A/N$$

$$\text{But } \text{var}(\hat{A}) \geq \frac{A}{N} \left(\frac{A}{A + 1/2} \right) < A/N$$

Even as $N \rightarrow \infty$, $\bar{x}$ does not attain CRLB. Thus, MLE is better (at least for large data records). For finite data records we would need to determine the exact mean and variance of $\hat{A}$ and compare them to $\bar{x}$.

$$\begin{aligned}
 3) \quad a) \quad p(\underline{x}; \mu) &= \prod_{n=0}^{N-1} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x[n] - \mu)^2} \\
 &= \frac{1}{(2\pi)^{N/2}} e^{-\frac{1}{2} \sum_{n=0}^{N-1} (x[n] - \mu)^2}
 \end{aligned}$$

To maximize p , we minimize $\sum_n (x[n] - \mu)^2$. Since it is a quadratic in μ , differentiation produces a global minimum.

$$\Rightarrow \sum_n (x[n] - \mu) = 0 \Rightarrow \hat{\mu} = \bar{x}$$

(This is just a DC level, μ , in WGN).

$$b) \quad p(\underline{x}; \lambda) = \lambda^N e^{-\lambda \sum_n x[n]} \quad \begin{array}{l} \text{all } x[n] > 0 \\ 0 \quad \text{otherwise} \end{array}$$

Assuming all $x[n] > 0$ we have

$$p = \lambda^N e^{-\lambda N \bar{x}}$$

$$\frac{dp}{d\lambda} = N \lambda^{N-1} e^{-\lambda N \bar{x}} + \lambda^N (-N \bar{x}) e^{-\lambda N \bar{x}} = 0$$

$$\Rightarrow \hat{\lambda} = 1/\bar{x}$$

To verify that $\hat{\lambda}$ yields the global maximum we can consider $\ln p$, which is a monotonic function of p .

$$\ln p = N \ln \lambda - \lambda N \bar{x}$$

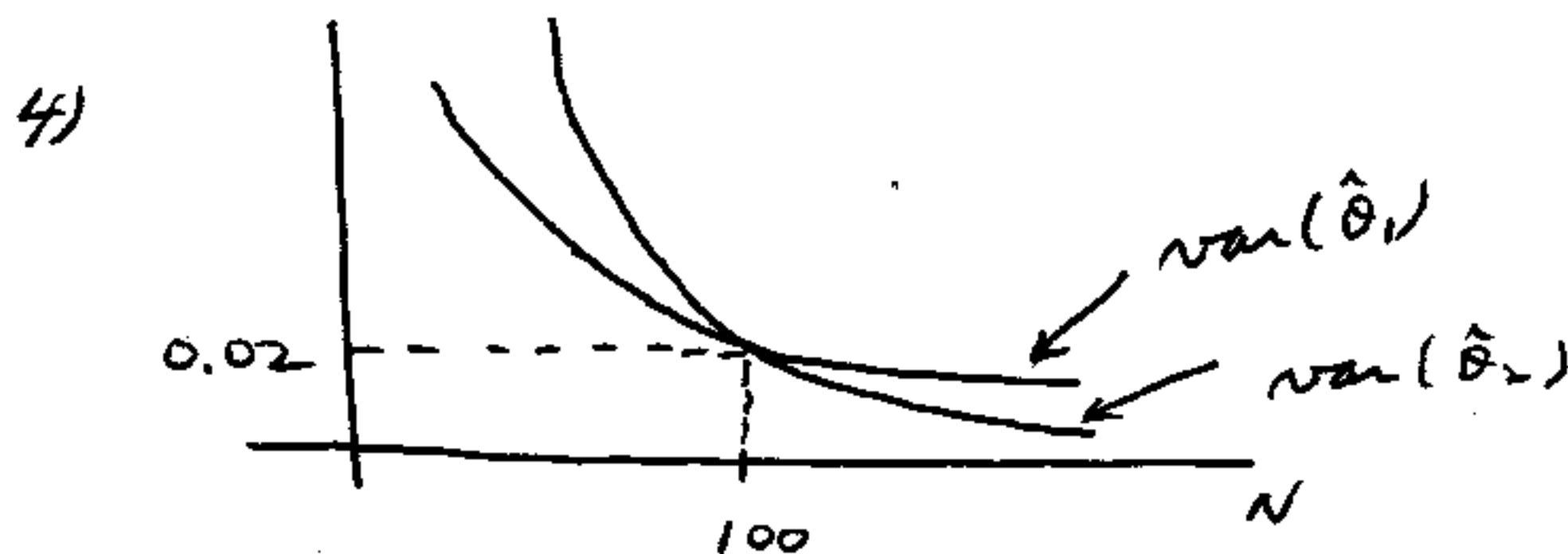
$$\frac{d \ln p}{d \lambda} = N/\lambda - N \bar{x}$$

$$\frac{d^2 \ln p}{d \lambda^2} = -N/\lambda^2 < 0 \quad \text{for all } \lambda$$

$\Rightarrow \ln p$ is concave function

$\Rightarrow \hat{\lambda}$ is global maximum solution

This result is reasonable since the mean of x is $1/\lambda$.



$\hat{\theta}_1$ better for $N < 100$
 $\hat{\theta}_2$ better for $N > 100$

$$5) \quad P\{|\hat{A} - A| > \epsilon\} = P\left\{\left|\frac{\bar{x} - A}{\sigma/\sqrt{N}}\right| > \frac{\epsilon}{\sigma/\sqrt{N}}\right\}$$

$$\leq \frac{1}{(\epsilon/\sigma/\sqrt{N})^2}$$

$$= \frac{\sigma^2}{N \epsilon^2} \rightarrow 0 \text{ as } N \rightarrow \infty$$

$\Rightarrow \bar{x}$ is consistent

- 6) Linearizing about the true value of θ , i.e., θ_0 yields

$$\alpha = g(\theta) = g(\theta_0) + \left. \frac{dg}{d\theta} \right|_{\theta=\theta_0} (\theta - \theta_0)$$

$$\begin{aligned} \text{But } \hat{\alpha} - \alpha &= g(\hat{\theta}) - g(\theta_0) \\ &\approx \left[g(\theta_0) + \left. \frac{dg}{d\theta} \right|_{\theta=\theta_0} (\hat{\theta} - \theta_0) \right] - g(\theta_0) \\ &= \left. \frac{dg}{d\theta} \right|_{\theta=\theta_0} (\hat{\theta} - \theta_0) \end{aligned}$$

$$\begin{aligned} P_r \{ |\hat{\alpha} - \alpha| > \epsilon \} &= P_r \left\{ \left| \left. \frac{dg}{d\theta} \right|_{\theta=\theta_0} (\hat{\theta} - \theta_0) \right| > \epsilon \right\} \\ &= P_r \left\{ |\hat{\theta} - \theta_0| > \frac{\epsilon}{\left| \left. \frac{dg}{d\theta} \right|_{\theta=\theta_0} \right|} \right\} \rightarrow 0 \text{ as } N \rightarrow \infty \end{aligned}$$

since $\hat{\theta}$ is consistent for θ (as long as $dg/d\theta$ is bounded).

$$\begin{aligned} 7) \quad p(\underline{x}; \theta) &= \prod_{n=0}^{N-1} e^{A(\theta)B(x(n)) + C(x(n)) + D(\theta)} \\ &= e^{A(\theta) \sum_n B(x(n)) + \sum_n C(x(n)) + ND(\theta)} \end{aligned}$$

To maximize p we must minimize

$$A(\theta) \sum_n B(x(n)) + ND(\theta)$$

Differentiating produces the necessary condition

$$\frac{dA(\theta)}{d\theta} \sum_{n=0}^N B(x(n)) + N \frac{dD(\theta)}{d\theta} = 0$$

For the PDFs of Prob. 7.2

$$a) p(x; \mu) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x^2 - 2\mu x + \mu^2)}$$

$$\Rightarrow A(\theta) = \mu \quad B(x) = x \quad D(\theta) = -\frac{1}{2} \mu^2$$

$$(i) \sum_n x(n) + N(-\mu) = 0 \Rightarrow \hat{\mu} = \bar{x}$$

$$b) p(x; \lambda) = \begin{cases} \lambda e^{-\lambda x} & x > 0 \\ 0 & x < 0 \end{cases}$$

$$= \begin{cases} e^{-\lambda x + \ln \lambda} & x > 0 \\ 0 & x < 0 \end{cases}$$

$$\Rightarrow A(\theta) = -\lambda \quad B(x) = x \quad D(\theta) = \ln \lambda$$

$$(-i) \sum_n x(n) + N(1/\lambda) = 0 \Rightarrow \hat{\lambda} = 1/\bar{x}$$

8) For discrete random variables the ML principle applies ^{to} the probability function

$$\begin{aligned} p_r\{x\} &= \prod_{n=0}^{N-1} p^{x(n)} (1-p)^{1-x(n)} \\ &= p^{\sum_n x(n)} (1-p)^{N - \sum_n x(n)} \end{aligned}$$

$$= p^{N\bar{x}} (1-p)^{N-N\bar{x}}$$

Maximizing $\ln P\{x\}$ over p

$$\frac{d \ln P\{x\}}{dp} = \frac{d}{dp} [N\bar{x} \ln p + (N-N\bar{x}) \ln (1-p)]$$

$$= \frac{N\bar{x}}{p} + \frac{N-N\bar{x}}{1-p} (-1) = 0$$

$$N\bar{x}(1-p) - (N-N\bar{x})p = 0 \Rightarrow \hat{p} = \bar{x}$$

$$9) \quad p(x) = \begin{cases} 1/\theta & 0 < x < \theta \\ 0 & \text{otherwise} \end{cases}$$

$$p(\underline{x}) = \prod_{n=0}^{N-1} p(x(n)) = \begin{cases} \frac{1}{\theta^N} & 0 < \text{all } x(n) < \theta \\ 0 & \text{otherwise} \end{cases}$$

Clearly, $p(\underline{x})$ is maximized over θ when θ is as small as possible. But $\theta > x(n)$ for all n . Thus, $\theta_{\min} = \max x(n)$ and thus $\hat{\theta} = \max x(n)$.

$$10) \quad p(\underline{x}; A) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_n (x(n) - AS(n))^2}$$

Minimize $\sum_n (x(n) - AS(n))^2$ to find MLE.

$$-2 \sum_n (x[n] - AS[n]) S[n] = 0$$

$$\Rightarrow \hat{A} = \frac{\sum_{n=0}^{N-1} x[n] S[n]}{\sum_{n=0}^{N-1} S^2[n]}$$

$$\begin{aligned} E(\hat{A}) &= \frac{\sum_n E(x[n]) S[n]}{\sum_n S^2[n]} = \frac{\sum_n AS[n] S[n]}{\sum_n S^2[n]} \\ &= A \end{aligned}$$

$$\begin{aligned} \text{var}(\hat{A}) &= \frac{\sum_n \text{var}(x[n]) S^2[n]}{(\sum_n S^2[n])^2} \\ &= \sigma^2 \frac{\sum_n S^2[n]}{(\sum_n S^2[n])^2} = \frac{\sigma^2}{\sum_{n=0}^{N-1} S^2[n]} \\ &= \mathbf{I}^{-1}(A) \quad (\text{see Theorem 4.1}) \end{aligned}$$

$\hat{A}$ is Gaussian since it is linear in the data. Hence, $\hat{A} \sim N(A, \mathbf{I}^{-1}(A))$ and thus asymptotic PDF holds for finite data records. This problem is special case of linear model.

$$\begin{aligned}
 \text{ii) } p(\underline{x}; \rho) &= \prod_{n=0}^{N-1} \frac{1}{2\pi \sqrt{\det(\underline{C})}} e^{-\frac{1}{2} \underline{x}^T(n) \underline{C}^{-1} \underline{x}(n)} \\
 &= \frac{1}{(2\pi)^N [\det(\underline{C})]^{N/2}} e^{-\frac{1}{2} \sum_n \underline{x}^T(n) \underline{C}^{-1} \underline{x}(n)}
 \end{aligned}$$

$$\text{But } \underline{C} = \begin{bmatrix} 1 & \rho \\ \rho & 1 \end{bmatrix} \Rightarrow \underline{C}^{-1} = \frac{\begin{bmatrix} 1 & -\rho \\ -\rho & 1 \end{bmatrix}}{1-\rho^2}$$

$$\det(\underline{C}) = 1-\rho^2$$

$$\begin{aligned}
 \ln p &= -N \ln 2\pi - \frac{N}{2} \ln(1-\rho^2) \\
 &\quad - \frac{1}{2(1-\rho^2)} \underbrace{\sum_n \underline{x}^T(n) \begin{bmatrix} 1 & -\rho \\ -\rho & 1 \end{bmatrix} \underline{x}(n)}_Q
 \end{aligned}$$

$$\begin{aligned}
 Q &= \sum_n (x_1(n) \ x_2(n)) \begin{bmatrix} 1 & -\rho \\ -\rho & 1 \end{bmatrix} \begin{bmatrix} x_1(n) \\ x_2(n) \end{bmatrix} \\
 &= \sum_n [x_1^2(n) + x_2^2(n) - 2\rho x_1(n) x_2(n)]
 \end{aligned}$$

$$\text{Let } C_{ij} = \sum_n x_i(n) x_j(n)$$

$$\Rightarrow Q = C_{11} + C_{22} - 2\rho C_{12}$$

$$\begin{aligned}
 \ln p &= -N \ln 2\pi - \frac{N}{2} \ln(1-\rho^2) \\
 &\quad - \frac{1}{2(1-\rho^2)} (C_{11} + C_{22} - 2\rho C_{12})
 \end{aligned}$$

$$\frac{d \ln p}{d\rho} = \frac{-N/2(-2\rho)}{1-\rho^2} - \frac{1}{2(1-\rho^2)} (-2C_{12})$$

$$- \frac{1}{2} (C_{11} + C_{22} - 2\rho C_{12}) \left[\frac{2\rho}{(1-\rho^2)^2} \right] = 0$$

$$N\rho(1-\rho^2) + C_{12}(1-\rho^2) - \rho(C_{11} + C_{22} - 2\rho C_{12}) = 0$$

$$\rho^3 - \frac{1}{N} C_{12} \rho^2 + \left(\frac{1}{N} C_{11} + \frac{1}{N} C_{22} - 1 \right) \rho - \frac{1}{N} C_{12} = 0$$

As $N \rightarrow \infty$, $\frac{C_{11}}{N} \rightarrow 1$, $\frac{C_{22}}{N} \rightarrow 1$

$$\rho^3 - \frac{1}{N} C_{12} \rho^2 + \rho - \frac{1}{N} C_{12} = 0$$

for which a solution is $\hat{\rho} = \frac{1}{N} C_{12}$
 $= \frac{1}{N} \sum_{n=0}^{N-1} x_1(n) x_2(n)$

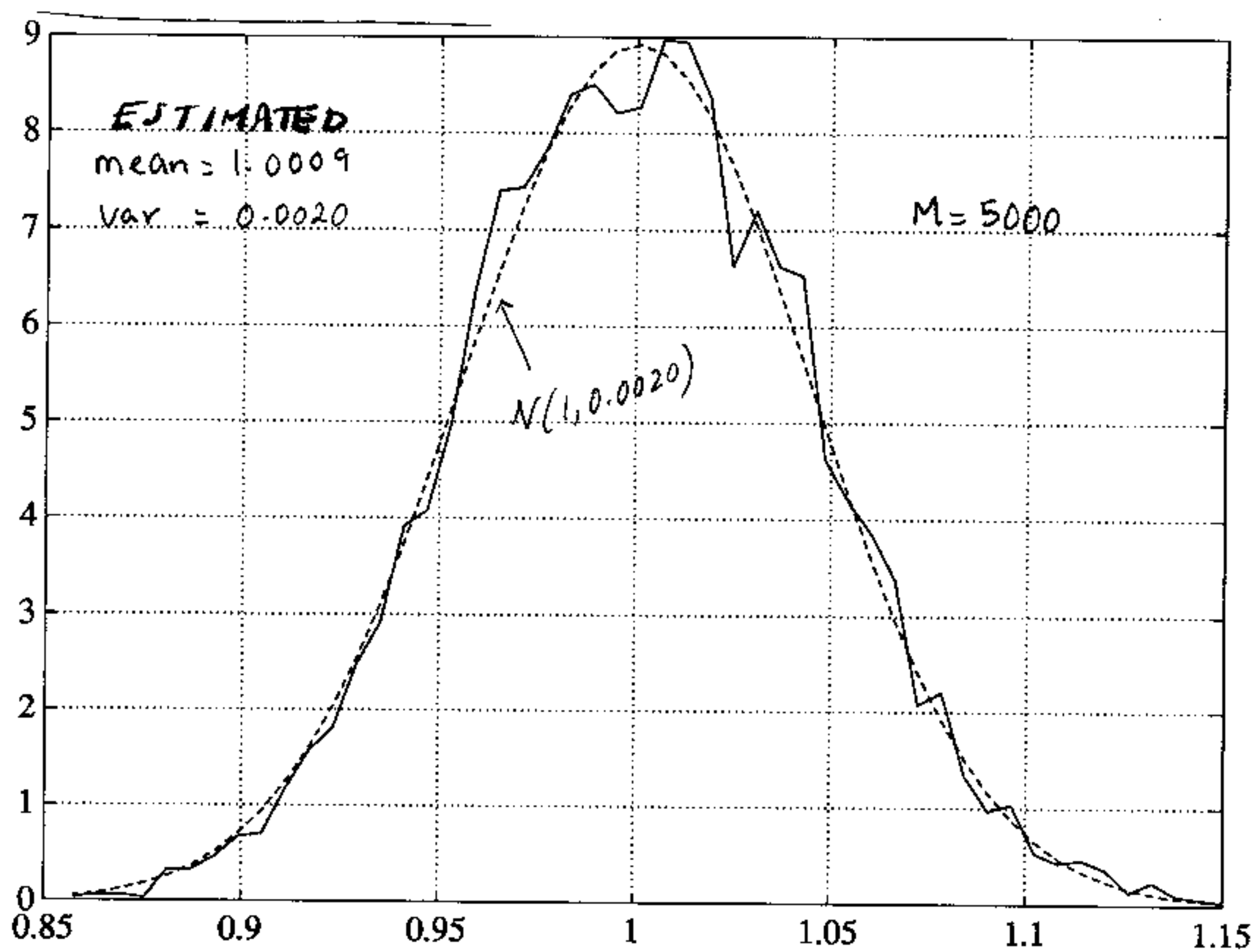
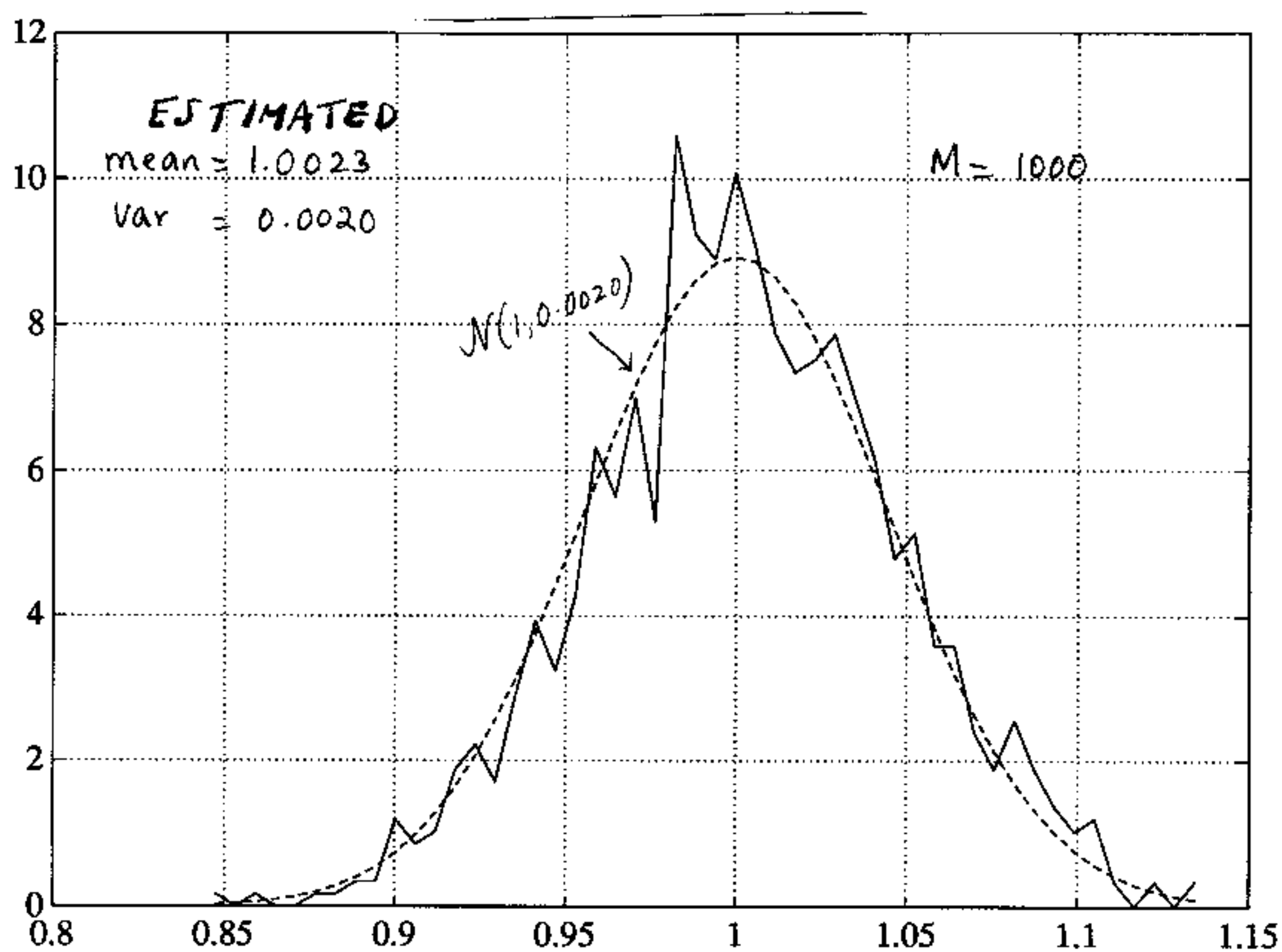
12) To find the MLE we maximize p or equivalently $\ln p$

$$\Rightarrow \frac{\partial \ln p}{\partial \theta} = 0$$

$$\text{But } \frac{\partial \ln p}{\partial \theta} = I(\theta) (\hat{\theta} - \theta)$$

Since $I(\theta) > 0$ for all θ (we always assume this - otherwise PDF does not depend on θ), the only solution is $\theta = \hat{\theta}$ or MLE is just $\hat{\theta}$, the efficient estimator.

13) See plots below.



We have a better fit as M increases.

14) See plots on next two pages.

15) E_s, E_c are zero mean since $W(n)$ is zero mean.

$$E(E_s E_c) = \frac{-4}{N^2 A^2} \sum_m \sum_n E(W(m) W(n)) \cdot \sin 2\pi f_0 m \cos 2\pi f_0 n$$

$$= \frac{-4}{N^2 A^2} \sigma^2 \sum_n \sin 2\pi f_0 n \cos 2\pi f_0 n$$

$$= \frac{-2\sigma^2}{N A^2} \sum_n \sin 4\pi f_0 n \approx 0$$

Thus, E_s, E_c are independent Gaussian random variables with zero means and variances given as follows:

$$\text{var}(E_s) = E(E_s^2) = \frac{4}{N^2 A^2} \sum_m \sum_n E(W(m) W(n))$$

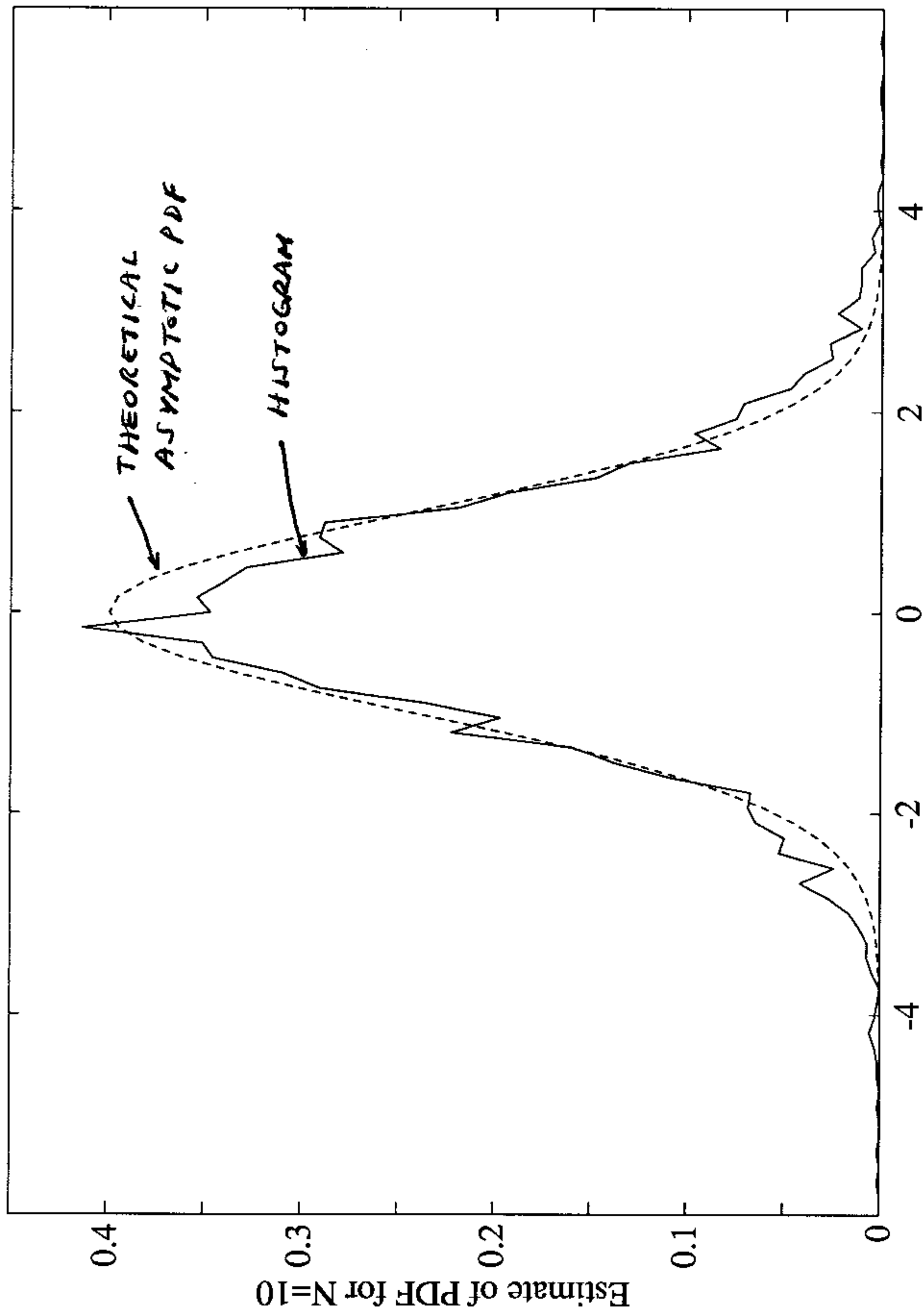
$$\cdot \sin 2\pi f_0 m \sin 2\pi f_0 n$$

$$= \frac{4\sigma^2}{N^2 A^2} \sum_n \sin^2 2\pi f_0 n$$

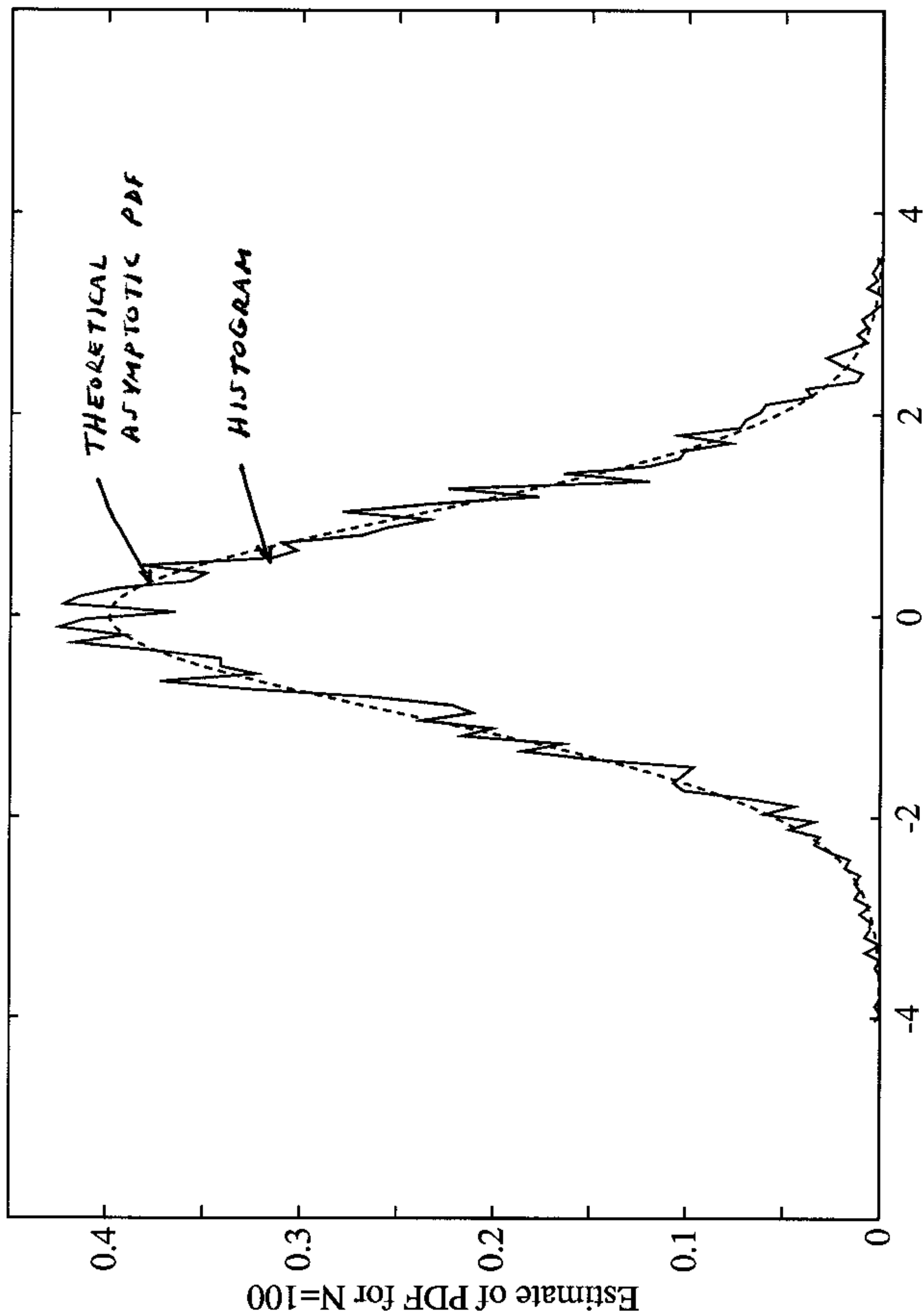
$$= \frac{4\sigma^2}{N^2 A^2} \sum_n \left(\frac{1}{2} - \frac{1}{2} \cos 4\pi f_0 n \right)$$

$$\approx \frac{4\sigma^2}{N^2 A^2} (N/2) = \frac{2\sigma^2}{N A^2}$$

Problem 7.14: Verification of Slutsky's theorem



Problem 7.14: Verification of Slutsky's theorem



and similarly, $\text{var}(\epsilon_c) \approx 2\sigma^2/NA^2$.

$$g(\epsilon_s, \epsilon_c) = \arctan \frac{\sin \phi + \epsilon_s}{\cos \phi + \epsilon_c}$$

$$\frac{\partial g}{\partial \epsilon_s} = \frac{1}{1 + \left(\frac{\sin \phi + \epsilon_s}{\cos \phi + \epsilon_c} \right)^2} \cdot \frac{1}{\cos \phi + \epsilon_c}$$

$$\left. \frac{\partial g}{\partial \epsilon_s} \right|_{\epsilon_s = \epsilon_c = 0} = \frac{\cos^2 \phi}{\cos^2 \phi + \sin^2 \phi} \cdot \frac{1}{\cos \phi} = \cos \phi$$

$$\frac{\partial g}{\partial \epsilon_c} = \frac{1}{1 + \left(\frac{\sin \phi + \epsilon_s}{\cos \phi + \epsilon_c} \right)^2} \cdot \frac{\sin \phi + \epsilon_s}{(\cos \phi + \epsilon_c)^2}$$

$$\left. \frac{\partial g}{\partial \epsilon_c} \right|_{\epsilon_s = \epsilon_c = 0} = \cos^2 \phi \frac{\sin \phi}{\cos^2 \phi} = \sin \phi$$

$$\hat{\phi} \approx \phi + \cos \phi \epsilon_s + \sin \phi \epsilon_c$$

$$\Rightarrow \hat{\phi} \sim N\left(\phi, \underbrace{\cos^2 \phi \text{var}(\epsilon_s) + \sin^2 \phi \text{var}(\epsilon_c)}_{\frac{2\sigma^2}{NA^2}}\right)$$

The variance is just the CRLB.

$$\begin{aligned} 16) \quad \text{var}(\hat{A}) &= \text{var}(x[0]) \\ &= \text{var}(W(0)) \\ &= \int_{-\infty}^{\infty} u^2 \frac{1}{2} e^{-|u|} du \end{aligned}$$

$$= \int_0^{\infty} u^2 e^{-u} du$$

$$= - (u^2 + 2u + 2) e^{-u} \Big|_0^{\infty} = 2$$

$$\text{var}(\hat{A}) \geq \int_{-\infty}^{\infty} \left(\frac{d p(u)}{du} \right)^2 / p(u) du$$

$$= \int_{-\infty}^{\infty} \left(\frac{1}{2} e^{-|u|} \right)^2 / \frac{1}{2} e^{-|u|} du$$

$$\text{since } dp/du = \begin{cases} -\frac{1}{2} e^{-u} & u > 0 \\ \frac{1}{2} e^u & u < 0 \end{cases}$$

$$\text{var}(\hat{A}) \geq \int_{-\infty}^{\infty} \frac{1}{2} e^{-|u|} du = 1$$

No, the MLE has variance 2 for all N .

17) Let $\alpha = 1/\theta$ so that by invariance of the MLE, $\hat{\alpha} = 1/\hat{\theta}$. But $\hat{\alpha} = 1/N \sum_{n=0}^{N-1} x^2(n)$

since

$$p(\underline{x}; \alpha) = \frac{1}{(2\pi\alpha)^{N/2}} e^{-\frac{1}{2\alpha} \sum_{n=0}^{N-1} x^2(n)}$$

$$\frac{\partial \log}{\partial \alpha} = -\frac{N}{2} \frac{1}{\alpha} + \frac{1}{2\alpha^2} \sum x^2(n) = 0$$

$$\Rightarrow \hat{\alpha} = \frac{1}{N} \sum_{n=0}^{N-1} x^2(n)$$

$$\text{Thus, } \hat{\theta} = \frac{1}{\frac{1}{N} \sum_{n=0}^{N-1} x^2(n)}$$

From Chapter 3

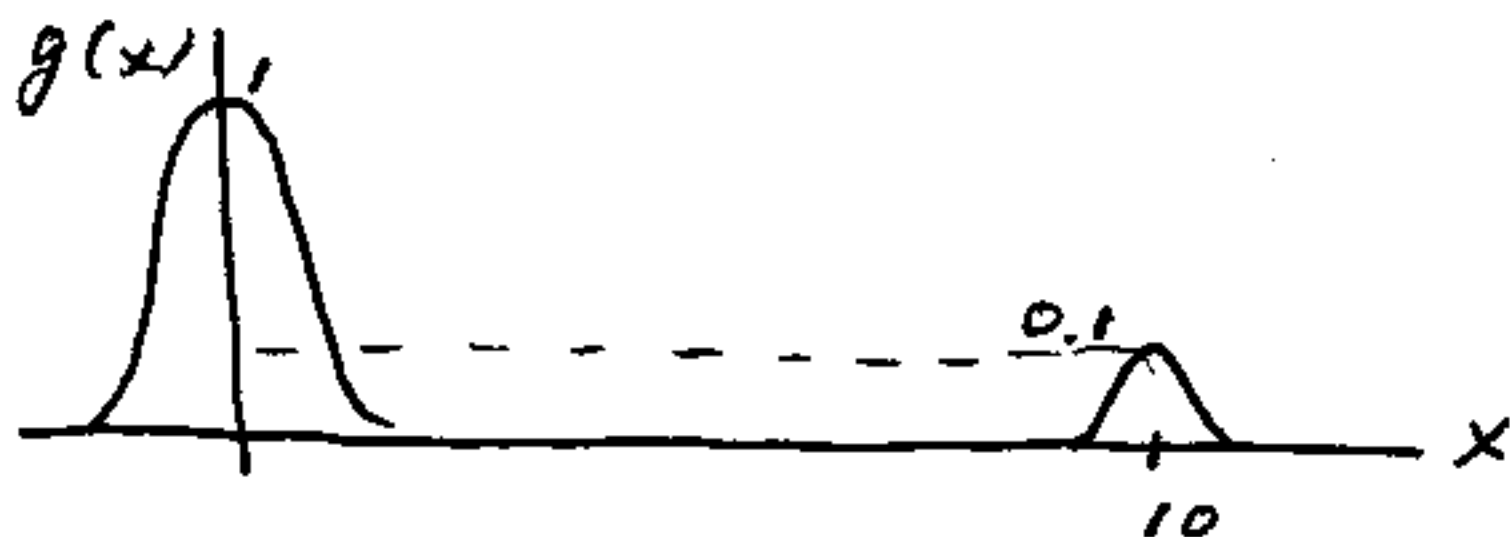
$$I(\alpha) = N/2\alpha^2$$

$$\text{and } I^{-1}(\alpha) = I^{-1}(\theta) (\partial \alpha / \partial \theta)^2$$

$$\begin{aligned} I^{-1}(\theta) &= \frac{2\alpha^2}{N} \left(\frac{\partial \theta}{\partial \alpha} \right)^2 \\ &= \frac{2\alpha^2}{N} (-1/\alpha^2)^2 = \frac{2}{N\alpha^2} \\ &= 2\theta^2/N \end{aligned}$$

$$\Rightarrow \hat{\theta} \sim N(\theta, 2\theta^2/N)$$

18)



$$\begin{aligned} g'(x) &= -x e^{-\frac{1}{2}x^2} - 0.1(x-10) e^{-\frac{1}{2}(x-10)^2} \\ g''(x) &= x^2 e^{-\frac{1}{2}x^2} + 0.1(x-10)^2 e^{-\frac{1}{2}(x-10)^2} \\ &\quad - 0.1 e^{-\frac{1}{2}(x-10)^2} - e^{-\frac{1}{2}x^2} \end{aligned}$$

$$x_{k+1} = x_k - \frac{dg/dx}{d^2g/dx^2} \Big|_{x=x_k}$$

Using a computer it is found that for $x_0 = 0.5$, the iteration converges to $x = 0$ and for $x_0 = 0.95$ it converges to $x = 10$. For other initial values such as $x_0 = 1$ it does not converge. To attain the maximum (global), x_0 must be close to the true value.

$$19) \quad p(\underline{x}; f_0) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_n (x[n] - \cos 2\pi f_0 n)^2}$$

To find MLE minimize

$$\begin{aligned} \sum_n (x[n] - \cos 2\pi f_0 n)^2 &= \\ &= \sum_n x^2[n] - 2 \sum_n x[n] \cos 2\pi f_0 n \\ &\quad + \sum_n \cos^2 2\pi f_0 n \end{aligned}$$

$$\approx \sum_n x^2[n] - 2 \sum_n x[n] \cos 2\pi f_0 n + N/2$$

$$\Rightarrow \text{maximize } \sum_{n=0}^{N-1} x[n] \cos 2\pi f_0 n$$

over $0 < f_0 < 1/2$

Using a Newton-Raphson iteration

$$g(f_0) = \sum_n x[n] \cos 2\pi f_0 n$$

$$dg/df_0 = - \sum_n 2\pi n x[n] \sin 2\pi f_0 n$$

$$d^2g/df_0^2 = - \sum_n (2\pi n)^2 x[n] \cos 2\pi f_0 n$$

$$f_{0,k+1} = f_{0,k} - \frac{\sum_n (2\pi n) x[n] \sin 2\pi f_{0,k} n}{\sum_n (2\pi n)^2 x[n] \cos 2\pi f_{0,k} n} \bigg|_{f_0 = f_{0,k}}$$

The function to be maximized is shown on following page for 10 different realizations of $w(n)$. Next, for 500 realizations we plot the results for a grid search and a Newton-Raphson iteration with 10 iterations. In (a) the grid search produces.

$$E(\hat{f}_0) = 0.25 = f_0$$

$$\text{var}(\hat{f}_0) = 1.3 \times 10^{-6}$$

In (b), (c), (d) the Newton-Raphson produces

Initial freq. (f_0)	$E(\hat{f}_0)$	$\text{var}(\hat{f}_0)$
0.24	0.25	1.2×10^{-6}
0.26	0.25	1.2×10^{-6}
0.28	0.16	10

$$20) \quad p(\underline{x}; \underline{\xi}) = \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_{n=0}^{N-1} (x[n] - \xi)^2}$$

Maximized when $\sum (x[n] - \xi)^2$ is minimized.

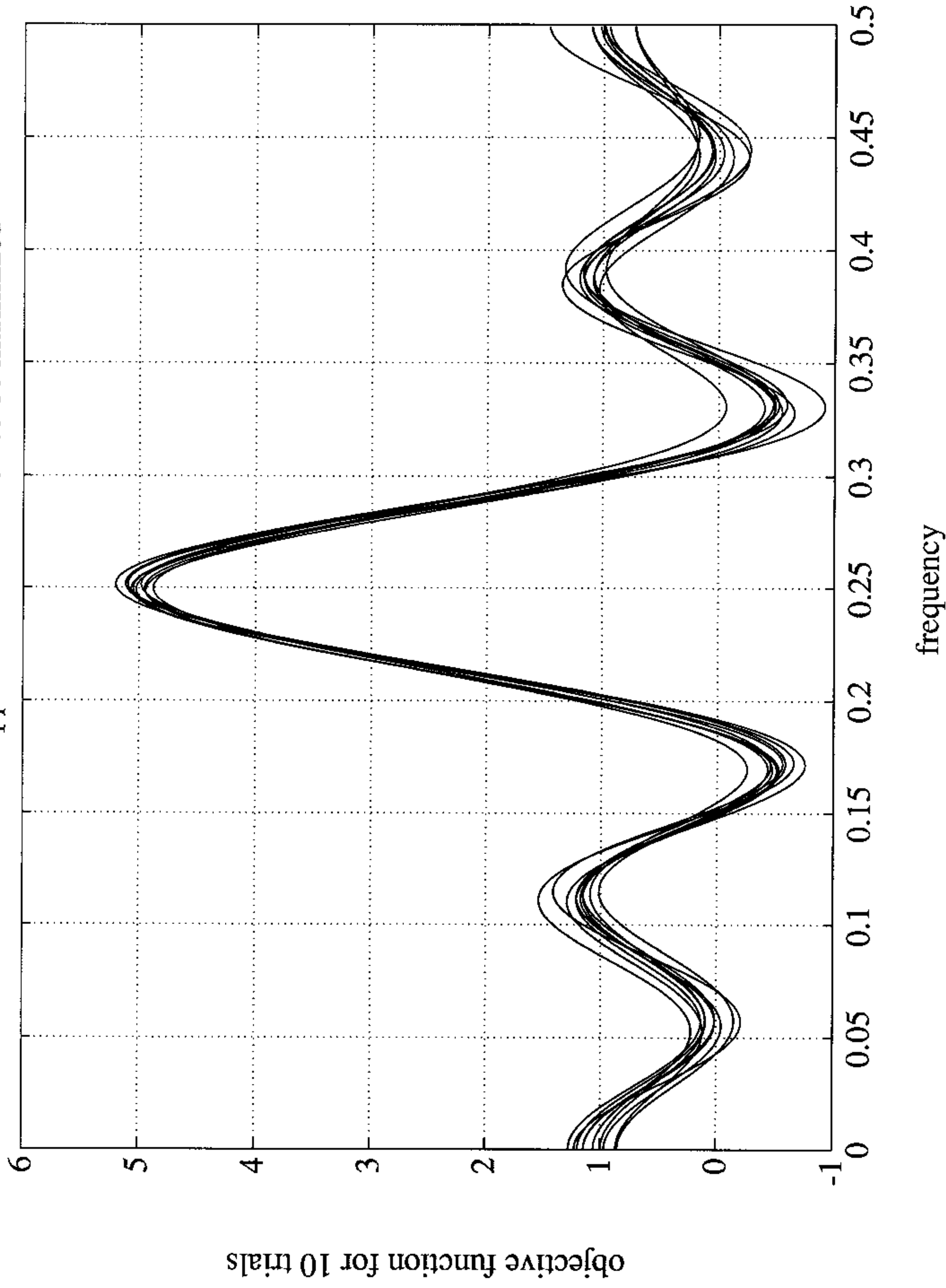
Clearly, then $\hat{\xi}[n] = x[n]$.

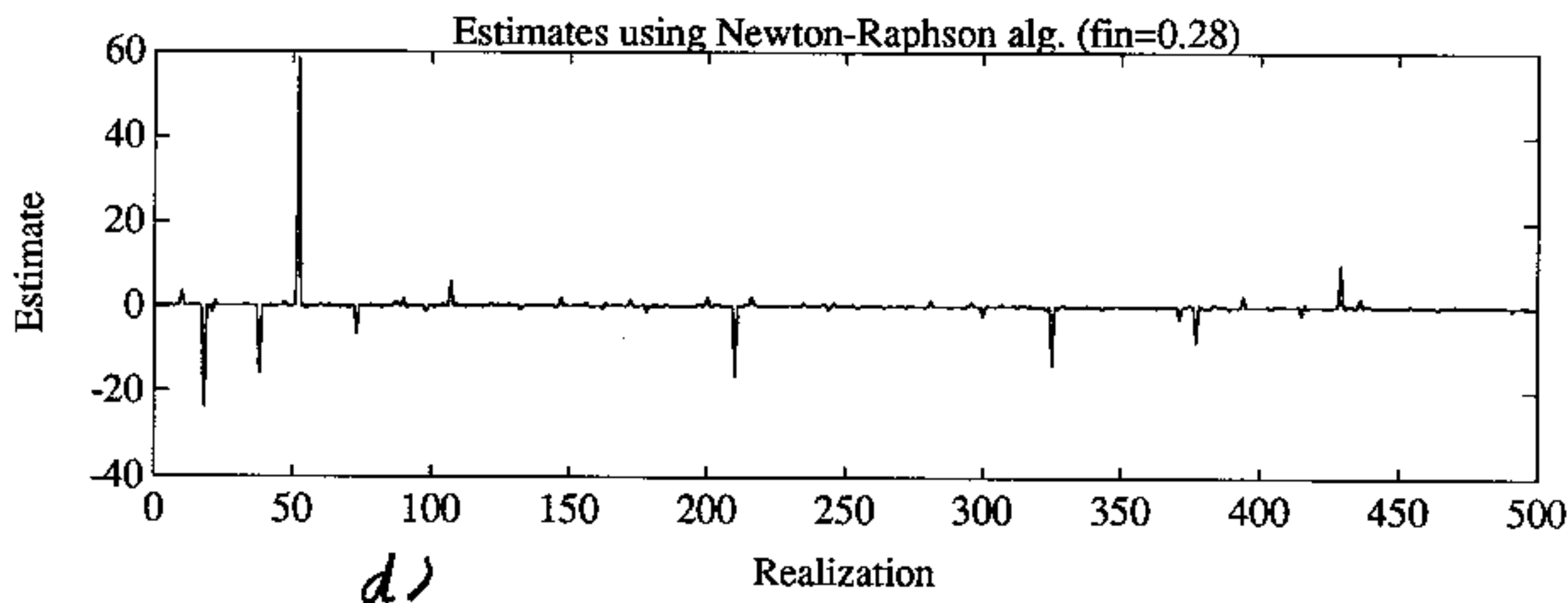
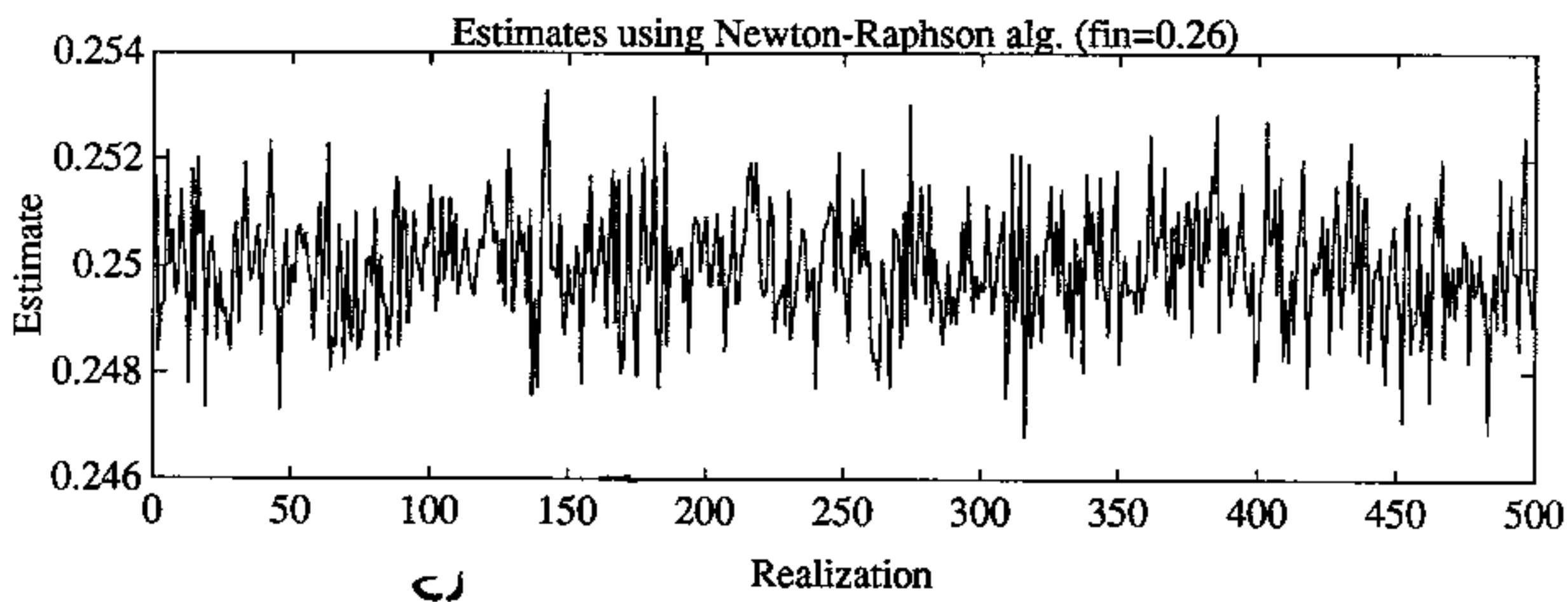
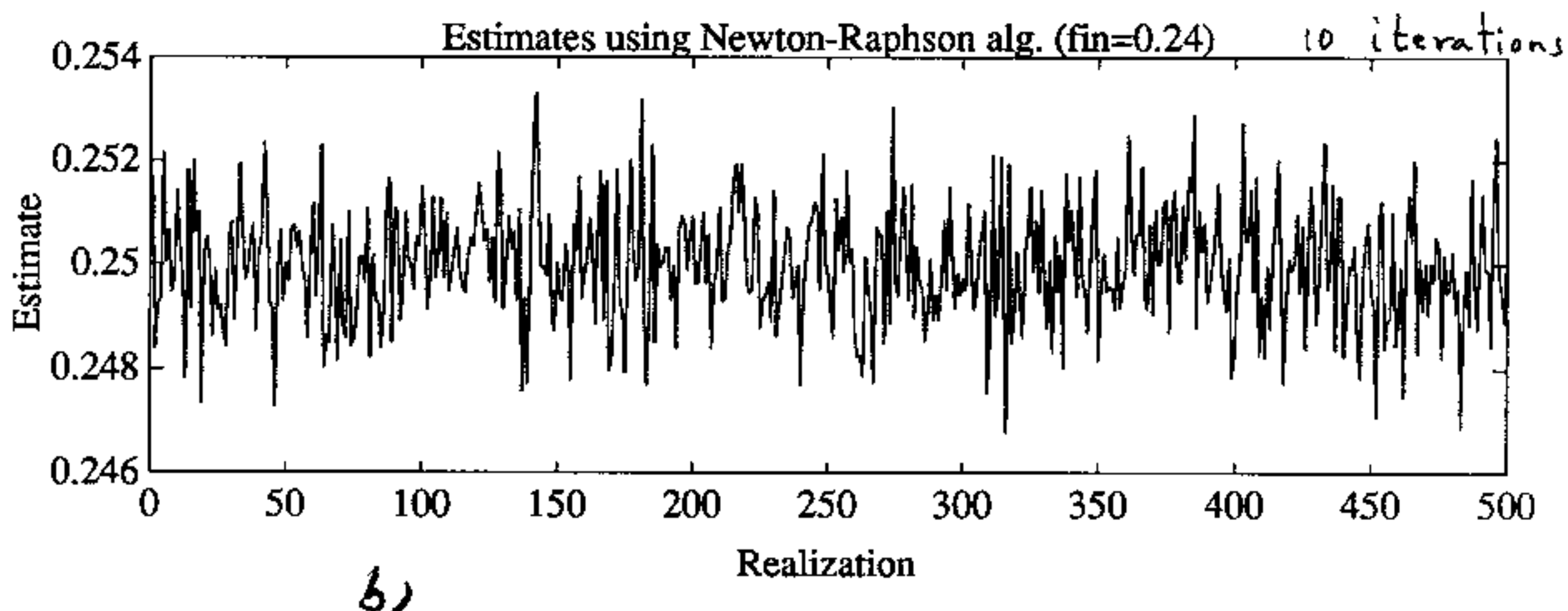
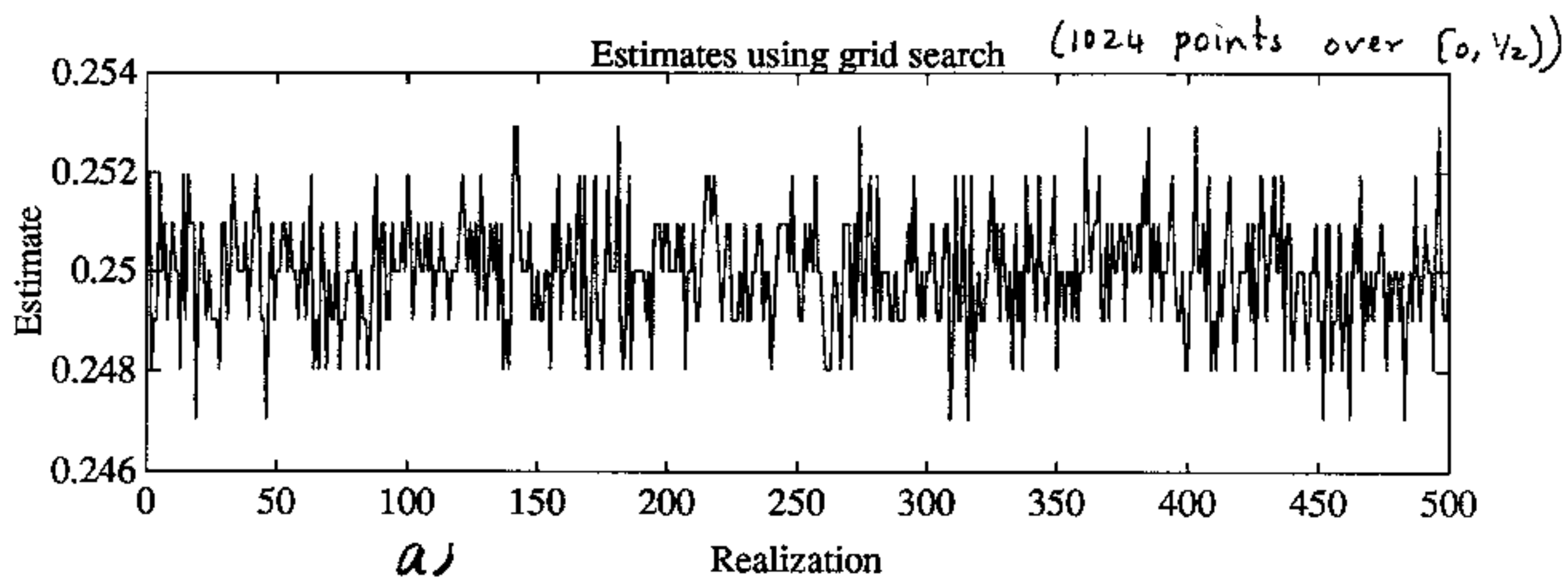
But $\hat{\xi} = \underline{x} \sim N(\underline{\xi}, \sigma^2 \underline{I})$ and thus $\hat{\xi}$ is unbiased and Gaussian. To determine if it is efficient we find the CRLB.

$$\frac{\partial \ln p}{\partial \xi[k]} = \frac{1}{\sigma^2} (x[k] - \xi[k])$$

$$\frac{\partial^2 \ln p}{\partial \xi[k] \partial \xi[l]} = -\frac{1}{\sigma^2} \delta_{kl}$$

Problem 7.19: Approx. MLE function to be maximized





$$\Rightarrow \underline{I}(\underline{\xi}) = \frac{1}{\sigma^2} \underline{I}$$

$$\text{or } \hat{\underline{\xi}} \sim N(\underline{\xi}, \underline{I}^{-1}(\underline{\xi}))$$

Hence, $\hat{\underline{\xi}}$ is efficient. However, it is not consistent since as $N \rightarrow \infty$, $\text{var}(\hat{\xi}[n]) = \sigma^2 \nrightarrow 0$.

21) From Example 7.12 $\hat{A} = \bar{X}$
 $\hat{\sigma}^2 = \frac{1}{N} \sum_{n=0}^{N-1} (X[n] - \bar{X})^2$

Hence, by the invariance

property $\hat{\xi} = \frac{\hat{A}^2}{\hat{\sigma}^2} = \frac{\bar{X}^2}{\frac{1}{N} \sum_{n=0}^{N-1} (X[n] - \bar{X})^2}$

$$22) I = \int_{-\frac{1}{2}}^{\frac{1}{2}} \ln |A(f)|^2 df = \int_{-\frac{1}{2}}^{\frac{1}{2}} [\ln A(f) + \ln A^*(f)] df$$

$$= 2 \operatorname{Re} \int_{-\frac{1}{2}}^{\frac{1}{2}} \ln A(f) df$$

Now let

$$z = e^{j2\pi f}, \quad dz = j2\pi e^{j2\pi f} df \\ = j2\pi z df$$

$$\Rightarrow I = 2 \operatorname{Re} \oint \ln A(z) \frac{1}{j2\pi} \frac{dz}{z}$$

$$= 2 \operatorname{Re} \left\{ \frac{1}{2\pi j} \oint \ln A(z) \frac{dz}{z} \right\}$$

$$= 2 \operatorname{Re} \left\{ \mathcal{Z}^{-1} \{ \ln A(z) \} \Big|_{n=0} \right\}$$

Since $A(z)$ converges for all $z \neq 0$, it

converges for $|z| \geq 1$ and since all of its zeros are within the unit circle, $\ln A(z)$ converges on and outside the unit circle. Thus, $\ln A(z)$ has a causal inverse.

The sample at $n=0$ is found from the initial value theorem or

$$\begin{aligned} \mathcal{Z}^{-1} \{ \ln A(z) \}_{n=0} &= \lim_{z \rightarrow \infty} \ln A(z) \\ &= \lim_{z \rightarrow \infty} \ln \left[1 + \sum_{k=1}^{\infty} a[k] z^{-k} \right] = \ln 1 \\ &= 0 \end{aligned}$$

2.3) We use (7.60) which when differentiated produces

$$\int_{-\frac{1}{2}}^{\frac{1}{2}} \left[\frac{1}{P_{xx}(f)} - \frac{I(f)}{P_{xx}^2(f)} \right] \frac{\partial P_{xx}(f)}{\partial P_0} df = 0$$

$$\int_{-\frac{1}{2}}^{\frac{1}{2}} \left[\frac{1}{P_0 Q(f)} - \frac{I(f)}{P_0^2 Q^2(f)} \right] Q(f) df = 0$$

$$\int_{-\frac{1}{2}}^{\frac{1}{2}} \left(P_0 - \frac{I(f)}{Q(f)} \right) df = 0$$

$$\Rightarrow \hat{P}_0 = \int_{-\frac{1}{2}}^{\frac{1}{2}} \frac{I(f)}{Q(f)} df$$

If $Q(f) = 1$ for all f ,

$$\hat{P}_0 = \int_{-\frac{1}{2}}^{\frac{1}{2}} I(f) df$$

$$= \frac{1}{N} \sum_{n=0}^{N-1} x^2[n] \quad \text{from Prob 7.25}$$

24) See plots on next page. For $f_0 = 0.25$ the peak of the periodogram and that of the exact function $\underline{x}^T \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{x}$ are at 0.25. But for $f_0 = 0.05$ the peak of the periodogram is shifted away ($= 0.068$) from the true value. This is due to the interaction of the complex sinusoids at $f_0 = 0.05$ and $-f_0 = -0.05$, which are not adequately resolved by the periodogram.

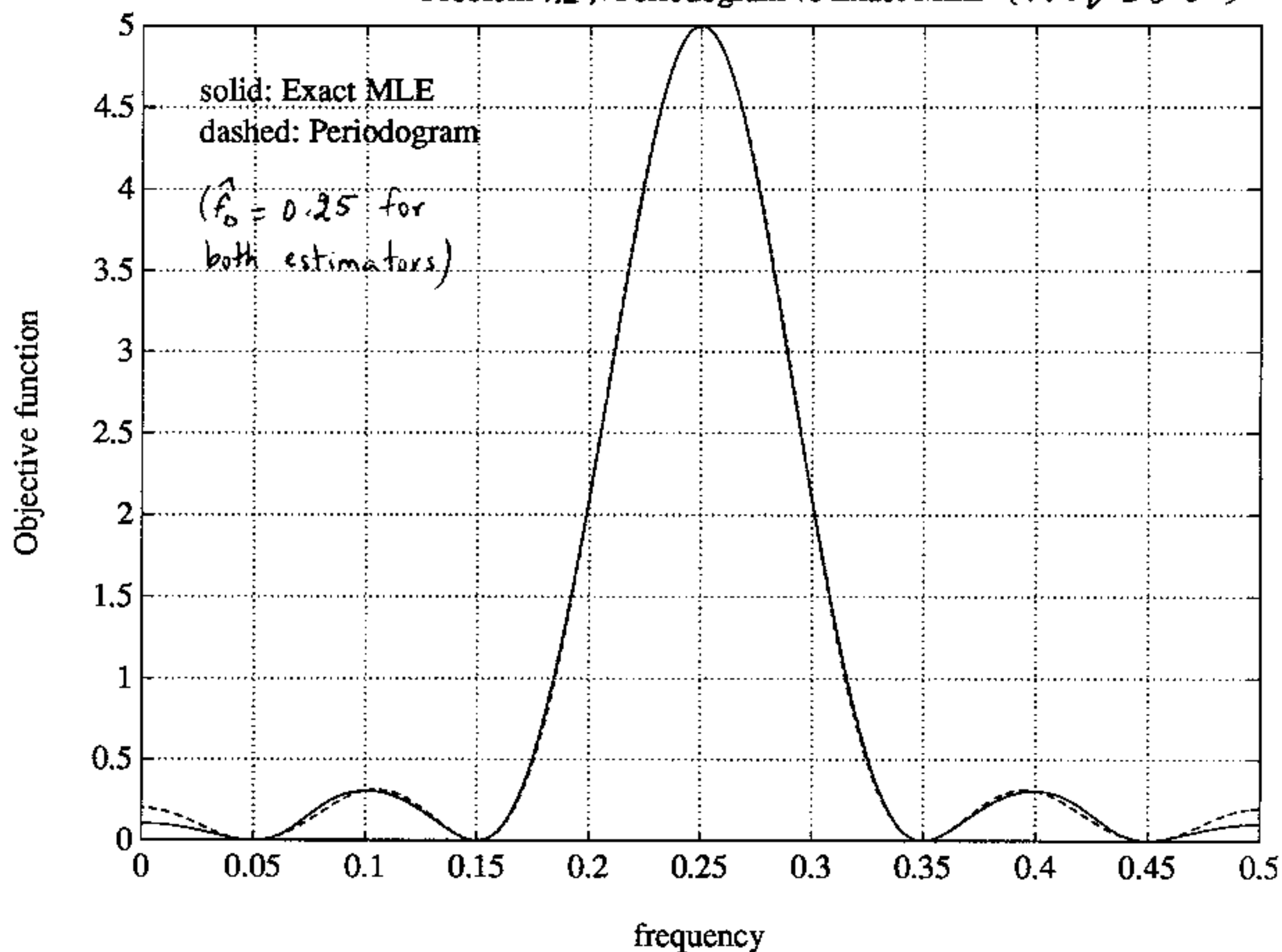
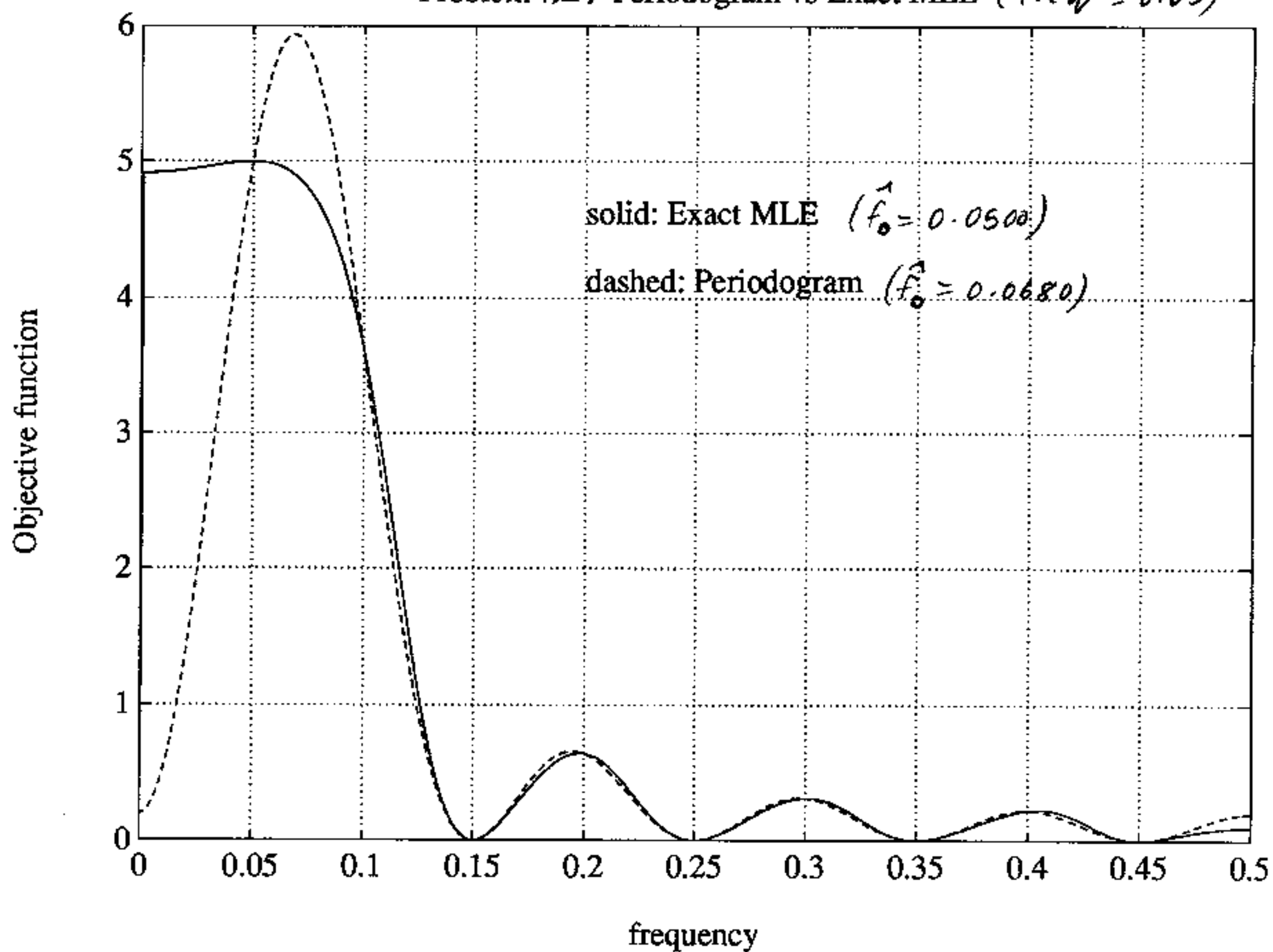
25) $\mathcal{F}^{-1} \{ I(f) \} = \frac{1}{N} \mathcal{F}^{-1} \{ x'(f) x'(f)^* \}$
 But if $x[n] \xleftrightarrow{\mathcal{F}} x'(f)$
 $x'[-n] \xleftrightarrow{\mathcal{F}} x'(f)^*$ for $x[n]$ real

$$\begin{aligned} \Rightarrow \mathcal{F}^{-1} \{ I(f) \} &= \frac{1}{N} x'[n] \star x'[-n] \\ &= \frac{1}{N} \sum_{n=-\infty}^{\infty} x'[-n] x'[k-n] \\ &= \frac{1}{N} \sum_{n=-\infty}^{\infty} x'[n] x'[n+k] \end{aligned}$$

For $k \geq 0$

$$\mathcal{F}^{-1} \{ I(f) \} = \frac{1}{N} \sum_{n=0}^{N-1-k} x[n] x[n+k]$$

since $x'[n] = 0$ for $n < 0$ or $n > N-1$
 $= x[n]$ otherwise

Problem 7.24 Periodogram vs Exact MLE ($f_{eq} = 0.25$)Problem 7.24 Periodogram vs Exact MLE ($f_{eq} = 0.05$)

For $k \geq 0$

$$\mathcal{I}^{-1} \{ \mathcal{I}(f) \} = \frac{1}{N} \sum_{n=-k}^{N-1} x[n] x[n+k]$$

Let $l = n+k$

$$= \frac{1}{N} \sum_{l=0}^{N-1+k} x[l-k] x[l]$$

$$= \frac{1}{N} \sum_{n=0}^{N-1+k} x[n] x[n-k]$$

Combining the results we have our solution.

$$\begin{aligned} 2b) \quad \hat{a}(1) &= -\hat{r}_{xx}[1] / \hat{r}_{xx}[0] \\ \hat{\sigma}^2 &= \hat{r}_{xx}[0] + \hat{a}(1) \hat{r}_{xx}[1] \end{aligned}$$

By invariance the MLE of $P_{xx}(f_0)$ is

$$\hat{P}_{xx}(f_0) = \frac{\hat{\sigma}^2}{|1 + \hat{a}(1) e^{-j2\pi f_0}|^2}$$

But

$$\hat{\mathcal{I}}_2 = \frac{\partial \mathcal{I}}{\partial \underline{\theta}} \mathcal{I}^{-1}(\underline{\theta}) \frac{\partial \mathcal{I}}{\partial \underline{\theta}}^T \geq 0$$

$$\underline{\mathcal{I}}(\underline{\theta}) = \begin{bmatrix} \frac{N \hat{r}_{xx}[0]}{\sigma_u^2} & 0 \\ 0 & \frac{N}{2\sigma_u^4} \end{bmatrix}$$

$$\text{And } g(a[1], \sigma_u^2) = \frac{\sigma_u^2}{|1 + a[1] e^{-j2\pi f_0}|^2}$$

so that

$$\partial g / \partial a[1] = \frac{-\sigma_u^2}{|A(f)|^4} (A(f) e^{j2\pi f_0} + A^*(f) e^{-j2\pi f_0})$$

$$= \frac{-\sigma_u^2}{|A(f)|^4} 2 \operatorname{Re}(A(f) e^{j2\pi f_0})$$

$$\partial g / \partial \sigma_u^2 = \frac{1}{|A(f)|^2}$$

$$\operatorname{var}(\hat{P}_{xx}(f_0)) \approx \frac{\partial g}{\partial \theta} \mathbf{I}^{-1}(\theta) \frac{\partial g}{\partial \theta}^T$$

$$= \frac{\sigma_u^2}{N r_{xx}[0]} \frac{\sigma_u^4}{|A(f)|^8} 4 \operatorname{Re}^2(A(f) e^{j2\pi f_0})$$

$$+ \frac{2\sigma_u^4}{N} \frac{1}{|A(f)|^4}$$

$$= \frac{4 P_{xx}^4(f_0)}{N r_{xx}[0] \sigma_u^2} \operatorname{Re}^2(A(f) e^{j2\pi f_0})$$

$$+ \frac{2}{N} P_{xx}^2(f_0)$$

Chapter 8

1) Nonlinear LS due to f_0 and r .

Yes, quadratic in A, B .

Analytically, we could find the values of A and B that minimize J for given f_0 and r . Then, plug these into J , which will now be a nonquadratic function of f_0 and r . Next, use a grid search over $0 < r < 1$ and $0 \leq f_0 \leq \frac{1}{2}$.

$$2) \quad J_{\min} = \sum_{n=0}^{N-1} x^2[n] - N \bar{x}^2 \leq \sum_{n=0}^{N-1} x^2[n]$$

$$\text{Also, } J_{\min} = \sum_{n=0}^{N-1} (x[n] - \bar{x})^2 \geq 0$$

$$3) \quad J = \sum_{n=0}^{M-1} (x[n] - A)^2 + \sum_{n=M}^{N-1} (x[n] + A)^2$$

$$\frac{\partial J}{\partial A} = -2 \sum_{n=0}^{M-1} (x[n] - A) + 2 \sum_{n=M}^{N-1} (x[n] + A) = 0$$

$$\Rightarrow -2 \sum_{n=0}^{M-1} x[n] + 2MA + 2 \sum_{n=M}^{N-1} x[n] + 2(N-M)A = 0$$

$$\hat{A} = \frac{1}{N} \left(\sum_{n=0}^{M-1} x[n] - \sum_{n=M}^{N-1} x[n] \right)$$

$$J_{\min} = \sum_{n=0}^{M-1} (x[n] - \hat{A})(x[n] - \hat{A}) + \sum_{n=M}^{N-1} (x[n] + \hat{A})(x[n] + \hat{A})$$

$$= \sum_{n=0}^{M-1} x[n](x[n] - \hat{A}) + \sum_{n=M}^{N-1} x[n](x[n] + \hat{A})$$

using the $\partial J / \partial A = 0$ equation.

$$\begin{aligned} J_{MIN} &= \sum_0^{N-1} x^2[n] - \hat{A} \left(\sum_0^{M-1} x[n] - \sum_M^{N-1} x[n] \right) \\ &= \sum_0^{N-1} x^2[n] - N \hat{A}^2 \end{aligned}$$

For $w[n]$ wgn we have

$$\begin{aligned} E(\hat{A}) &= 1/N [MA - (N-M)A] = A \\ \text{var}(\hat{A}) &= \frac{1}{N^2} \left(\text{var} \left(\sum_0^{M-1} x[n] \right) \right. \\ &\quad \left. + \text{var} \left(\sum_M^{N-1} x[n] \right) \right) \\ &= \frac{1}{N^2} \left(\sum_0^{M-1} \text{var}(x[n]) + \sum_M^{N-1} \text{var}(x[n]) \right) \\ &= \frac{1}{N^2} (M\sigma^2 + (N-M)\sigma^2) = \sigma^2/N \end{aligned}$$

$\Rightarrow \hat{A} \sim N(A, \sigma^2/N)$ since $\hat{A}$ is a linear function of the $x[n]$'s.

4) See solution for Prob 5.19

$$5) \quad \underline{z} = \underbrace{\begin{bmatrix} 1 & 1 & \dots & 1 \\ \cos 2\pi f_1 & \cos 2\pi f_2 & \dots & \cos 2\pi f_p \\ \vdots \\ \cos 2\pi f_1(N-1) & \cos 2\pi f_2(N-1) & \dots & \cos 2\pi f_p(N-1) \end{bmatrix}}_{\underline{H}} \underbrace{\begin{bmatrix} A_1 \\ A_2 \\ \vdots \\ A_p \end{bmatrix}}_{\underline{\theta}}$$

$$\underline{H}^T \underline{H} \hat{\underline{\theta}} = \underline{H}^T \underline{x} \quad \text{are normal equations}$$

If $f_i = i/N$, the column vectors of $\underline{H}$ are orthogonal (see (4.13)). Thus,

$$\underline{H}^T \underline{H} = \frac{N}{2} \underline{I}$$

$$\Rightarrow \hat{\underline{\theta}} = \frac{2}{N} \underline{H}^T \underline{x} \Rightarrow \hat{A}_i = \frac{2}{N} \sum_{n=0}^{N-1} x[n] \cos 2\pi f_i n$$

$$J_{M,N} = \underline{x}^T (\underline{I} - \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T) \underline{x}$$

$$= \underline{x}^T (\underline{I} - \frac{2}{N} \underline{H} \underline{H}^T) \underline{x}$$

$$= \underline{x}^T \underline{x} - \frac{2}{N} \|\underline{H}^T \underline{x}\|^2$$

$$= \underline{x}^T \underline{x} - \frac{2}{N} \left(\frac{N}{2}\right)^2 \|\hat{\underline{\theta}}\|^2$$

$$= \sum_{n=0}^{N-1} x^2[n] - \frac{N}{2} \sum_{i=1}^p \hat{A}_i^2$$

For $w[n]$, $w \in N$ the PDF is

$$\hat{\underline{\theta}} \sim N(\underline{\theta}, \sigma^2 (\underline{H}^T \underline{H})^{-1}) \quad \text{since}$$

$$E(\hat{\underline{\theta}}) = (\underline{H}^T \underline{H})^{-1} \underline{H}^T E(\underline{x}) = (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{H} \underline{\theta} = \underline{\theta}$$

$$C_{\hat{\underline{\theta}}} = E[(\hat{\underline{\theta}} - \underline{\theta})(\hat{\underline{\theta}} - \underline{\theta})^T]$$

$$= E\left[(\underline{H}^T \underline{H})^{-1} \underline{H}^T \underbrace{(\underline{x} - \underline{H} \underline{\theta})}_{\underline{w}} \underbrace{(\underline{x} - \underline{H} \underline{\theta})^T}_{\underline{w}^T} \underline{H} (\underline{H}^T \underline{H})^{-1} \right]$$

$$= (\underline{H}^T \underline{H})^{-1} \underline{H}^T \sigma^2 \underline{I} \underline{H} (\underline{H}^T \underline{H})^{-1} = \sigma^2 (\underline{H}^T \underline{H})^{-1}$$

$$\text{or } C_{\hat{\underline{\theta}}} = 2\sigma^2/N \underline{I}$$

and since $\hat{\underline{\theta}}$ is a linear function of $\underline{x}$,
we have a Gaussian PDF on

$$\hat{\underline{\theta}} \sim N(\underline{\theta}, 2\sigma^2/N \underline{I})$$

$$6) \quad \hat{\underline{\theta}} = (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{x}$$

From Prob 8.5 we have $\hat{\underline{\theta}} \sim N(\underline{\theta}, \sigma^2 (\underline{H}^T \underline{H})^{-1})$

Yes, it is unbiased.

$$\begin{aligned} 7) \quad E(\hat{\sigma}^2) &= \frac{1}{N} E \left[\underline{x}^T (\underline{I} - \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T) \underline{x} \right] \\ &= \frac{1}{N} \text{tr} \left[(\underline{I} - \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T) E(\underline{x} \underline{x}^T) \right] \end{aligned}$$

$$\text{Since } E(\underline{x} \underline{x}^T) = \text{tr} (E(\underline{y} \underline{x}^T))$$

$$E(\hat{\sigma}^2) = \frac{1}{N} \text{tr} \left[(\underline{I} - \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T) E((\underline{H}\underline{\theta} + \underline{w})(\underline{H}\underline{\theta} + \underline{w})^T) \right]$$

$$= \frac{1}{N} \text{tr} \left[(\underline{I} - \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T) (\underline{H} \underline{\theta} \underline{\theta}^T \underline{H}^T + \sigma^2 \underline{I}) \right]$$

$$\begin{aligned} &= \frac{1}{N} \text{tr} \left[\underline{H} \underline{\theta} \underline{\theta}^T \underline{H}^T + \sigma^2 \underline{I} \right. \\ &\quad \left. - \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{H} \underline{\theta} \underline{\theta}^T \underline{H}^T \right. \\ &\quad \left. - \sigma^2 \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \right] \end{aligned}$$

$$\begin{aligned}
&= \frac{\sigma^2}{N} \text{tr} \left\{ \underline{I} - \underline{H}(\underline{H}^T \underline{H})^{-1} \underline{H}^T \right\} \\
&= \sigma^2 - \frac{\sigma^2}{N} \text{tr} \left\{ \underline{H}(\underline{H}^T \underline{H})^{-1} \underline{H}^T \right\} \\
&= \sigma^2 - \sigma^2/N \underbrace{\text{tr} \left\{ (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{H} \right\}}_{\underline{I} \text{ (p \times p)}} \\
&= \sigma^2 - \sigma^2 p/N = \sigma^2 \frac{N-p}{N}
\end{aligned}$$

$\hat{\sigma}^2$ is biased. To make it unbiased use

$$\hat{\sigma}^2 = \frac{1}{N-p} J_{MIN}.$$

It is said that we lose p degrees of freedom in estimating $\underline{\theta}$.

$$\begin{aligned}
\text{Ex) } J(A) &= \sum_{n=0}^{N-1} \frac{1}{\sigma_n^2} (x(n) - A)^2 \\
\frac{\partial J}{\partial A} &= -2 \sum_n \frac{1}{\sigma_n^2} (x(n) - A) = 0 \\
\Rightarrow \hat{A} &= \frac{\sum_{n=0}^{N-1} x(n) / \sigma_n^2}{\sum_{n=0}^{N-1} \frac{1}{\sigma_n^2}}
\end{aligned}$$

$$E(\hat{A}) = \frac{\sum_n E(x(n)) / \sigma_n^2}{\sum_n 1 / \sigma_n^2} = A$$

$$\begin{aligned}
 \text{var}(\hat{A}) &= \frac{1}{\left(\sum_n 1/\sigma_n^2\right)^2} \underbrace{\sum_n \text{var}(x(n)/\sigma_n^2)}_{= \sum_n 1/\sigma_n^4 \text{var}(x(n))} \\
 &= \sum_n 1/\sigma_n^4 \text{var}(x(n)) \\
 &= \sum_n 1/\sigma_n^2 \\
 &= \frac{1}{\sum_{n=0}^{N-1} 1/\sigma_n^2}
 \end{aligned}$$

$$\begin{aligned}
 9) \quad J &= (\underline{X} - \underline{H}\underline{\theta})^T \underline{W} (\underline{X} - \underline{H}\underline{\theta}) \\
 &= (\underline{X} - \underline{H}\underline{\theta})^T \underline{D}^T \underline{D} (\underline{X} - \underline{H}\underline{\theta}) \\
 &= \underbrace{(\underline{D}\underline{X} - \underline{D}\underline{H}\underline{\theta})^T}_{\underline{X}' \quad \underline{H}'} (\underline{D}\underline{X} - \underline{D}\underline{H}\underline{\theta})
 \end{aligned}$$

$$\begin{aligned}
 \Rightarrow \hat{\underline{\theta}} &= (\underline{H}'^T \underline{H}')^{-1} \underline{H}'^T \underline{X}' \\
 &= (\underline{H}^T \underline{D}^T \underline{D} \underline{H})^{-1} \underline{H}^T \underline{D}^T \underline{D} \underline{X} \\
 &= (\underline{H}^T \underline{W} \underline{H})^{-1} \underline{H}^T \underline{W} \underline{X}
 \end{aligned}$$

$$\begin{aligned}
 J_{\min} &= \underline{X}'^T (\underline{I} - \underline{H}' (\underline{H}'^T \underline{H}')^{-1} \underline{H}'^T) \underline{X}' \\
 &= \underline{X}^T \underline{D}^T (\underline{I} - \underline{D} \underline{H} (\underline{H}^T \underline{D}^T \underline{D} \underline{H})^{-1} \underline{H}^T \underline{D}^T) \underline{D} \underline{X} \\
 &= \underline{X}^T (\underline{W} - \underline{W} \underline{H} (\underline{H}^T \underline{W} \underline{H})^{-1} \underline{H}^T \underline{W}) \underline{X}
 \end{aligned}$$

$$10) \quad \hat{\underline{y}} = \underline{H} \hat{\underline{\theta}} = \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{X} = \underline{P} \underline{X}$$

$$\underline{x} - \hat{\underline{s}} = \underline{x} - \underline{P}\underline{x} = (\underline{I} - \underline{P})\underline{x}$$

$$\begin{aligned} \|\hat{\underline{s}}\|^2 + \|\underline{x} - \hat{\underline{s}}\|^2 &= \underline{x}^T \underline{P}^T \underline{P} \underline{x} \\ &\quad + \underline{x}^T (\underline{I} - \underline{P})^T (\underline{I} - \underline{P}) \underline{x} \\ &= \underline{x}^T \underline{P} \underline{x} + \underline{x}^T (\underline{I} - \underline{P}) \underline{x} = \underline{x}^T \underline{x} = \|\underline{x}\|^2 \end{aligned}$$

Since $\underline{P}$ and $(\underline{I} - \underline{P})$ are symmetric and idempotent.

$$11) \quad \underline{x}_1 = \underline{z}_1 + \underline{z}_1^\perp \quad \underline{x}_2 = \underline{z}_2 + \underline{z}_2^\perp$$

$$\begin{aligned} l &= \underline{x}_1^T \underline{P} \underline{x}_2 - \underline{x}_2^T \underline{P} \underline{x}_1 = (\underline{z}_1 + \underline{z}_1^\perp)^T \underline{P} (\underline{z}_2 + \underline{z}_2^\perp) \\ &\quad - (\underline{z}_2 + \underline{z}_2^\perp)^T \underline{P} (\underline{z}_1 + \underline{z}_1^\perp) \\ &= (\underline{z}_1 + \underline{z}_1^\perp)^T \underline{P} \underline{z}_2 - (\underline{z}_2 + \underline{z}_2^\perp)^T \underline{P} \underline{z}_1 \end{aligned}$$

$$\text{Since } \underline{P} \underline{z}_1^\perp = 0$$

$$l = (\underline{z}_1 + \underline{z}_1^\perp)^T \underline{z}_2 - (\underline{z}_2 + \underline{z}_2^\perp)^T \underline{z}_1$$

Since $\underline{P} \underline{z}_2 = \underline{z}_2$

$$l = \underline{z}_1^T \underline{z}_2 - \underline{z}_2^T \underline{z}_1 = 0 \quad \text{Since } \underline{z}_i^T \underline{z}_j = 0$$

$$\text{Now let } \underline{x}_1 = \underline{e}_i = [0 \dots 0 \underset{\substack{\uparrow \\ i^{\text{th}} \text{ place}}}{1} 0 \dots 0]^T$$

$$\underline{x}_2 = \underline{e}_j$$

$$\underline{e}_i^T \underline{P} \underline{e}_j = \underline{e}_j^T \underline{P} \underline{e}_i \quad \text{from previous result}$$

$$\Rightarrow [\underline{P}]_{ij} = [\underline{P}]_{ji}$$

$$\begin{aligned} 12) a) \quad \underline{P}^2 &= \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T \\ &= \underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T = \underline{P} \end{aligned}$$

$$\begin{aligned} b) \quad \underline{x}^T \underline{P} \underline{x} &= \underline{x}^T \underline{P} \underline{P} \underline{x} = \underline{x}^T \underline{P}^T \underline{P} \underline{x} \\ &= (\underline{P} \underline{x})^T \underline{P} \underline{x} = \|\underline{P} \underline{x}\|^2 \geq 0 \end{aligned}$$

$$\begin{aligned} c) \quad \underline{P} \underline{x} &= \lambda \underline{x} \Rightarrow \underline{P} \underline{x} = \underline{P} \underline{P} \underline{x} = \underline{P} \lambda \underline{x} \\ &= \lambda (\underline{P} \underline{x}) = \lambda^2 \underline{x} \end{aligned}$$

$\Rightarrow$ If λ is an eigenvalue, so is λ^2 .
By uniqueness $\lambda^2 = \lambda$ or $\lambda = 0, 1$

d) Rank = sum of nonzero eigenvalues.
To show that there are p nonzero eigenvalues, we find $\text{tr}(\underline{P}) = \sum_{i=1}^N \lambda_i$

$$\begin{aligned} \text{tr}(\underline{P}) &= \text{tr}(\underline{H} (\underline{H}^T \underline{H})^{-1} \underline{H}^T) \\ &= \text{tr}((\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{H}) \\ &= \text{tr}(\underline{I}) = p \end{aligned}$$

$$\begin{aligned} \Rightarrow \sum_{i=1}^p \lambda_i &= p \Rightarrow p \text{ eigenvalues} = 1 \\ &\Rightarrow \text{rank} = p \end{aligned}$$

$$13) \quad \hat{\underline{\theta}} = (\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{x}$$

$$\text{where } \underline{H} = \begin{bmatrix} 1 & 0 \\ \vdots & \vdots \\ 1 & N-1 \end{bmatrix}$$

$$\underline{H}^T \underline{H} = \begin{bmatrix} N & \sum_n n \\ \sum_n n & \sum_n n^2 \end{bmatrix} = \begin{bmatrix} N & N(N-1)/2 \\ N(N-1)/2 & \frac{N(N-1)(2N-1)}{6} \end{bmatrix}$$

using (3.22)

$$(\underline{H}^T \underline{H})^{-1} = \frac{\begin{bmatrix} \frac{N(N-1)(2N-1)}{6} & -\frac{N(N-1)}{2} \\ -\frac{N(N-1)}{2} & N \end{bmatrix}}{\frac{N^2(N-1)(2N-1)}{6} - \frac{N^2(N-1)^2}{4}}$$

$$\hat{\underline{\theta}} = \begin{bmatrix} \frac{2(2N-1)}{N(N+1)} & \frac{-6}{N(N+1)} \\ \frac{-6}{N(N+1)} & \frac{12}{N(N^2-1)} \end{bmatrix} \begin{bmatrix} \sum_n x(n) \\ \sum_n n x(n) \end{bmatrix}$$

which produces (8.23).

$$14) \quad \underline{h}'_{k+1} \perp \{ \underline{h}_1, \underline{h}_2, \dots, \underline{h}_k \}$$

$$\text{Let } P_k = \underline{H}_k (\underline{H}_k^T \underline{H}_k)^{-1} \underline{H}_k^T$$

$$\text{where } \underline{H}_k = [\underline{h}_1 \dots \underline{h}_k]$$

$$\text{Then, } \underline{h}'_{k+1} = \underline{h}_{k+1} - P_k \underline{h}_{k+1} \quad (\text{See Fig 8.8}) \\ = P_k^\perp \underline{h}_{k+1}$$

Since $\underline{h}'_{k+1} \perp \{\underline{h}_1, \dots, \underline{h}_k\}$ it is also $\perp$ to $\hat{S}_k \Rightarrow$ we can project $\underline{x}$ onto $\underline{h}'_{k+1}$ and add the result to $\hat{S}_k$.

$$\alpha \underline{h}'_{k+1} = \frac{\underline{x}^T \underline{h}'_{k+1}}{\|\underline{h}'_{k+1}\|} \frac{\underline{h}'_{k+1}}{\|\underline{h}'_{k+1}\|}$$

Since $\frac{\underline{h}'_{k+1}}{\|\underline{h}'_{k+1}\|}$ is a unit vector in the $\underline{h}'_{k+1}$ direction

$$\text{or } \alpha = \frac{\underline{x}^T \underline{h}'_{k+1}}{\|\underline{h}'_{k+1}\|^2} = \frac{\underline{x}^T \underline{P}_k^\perp \underline{h}_{k+1}}{\|\underline{P}_k^\perp \underline{h}_{k+1}\|^2}$$

$$\hat{S}_{k+1} = \underline{H}_k \hat{\underline{\theta}}_k + \frac{\underline{x}^T \underline{P}_k^\perp \underline{h}_{k+1}}{\underline{h}_{k+1}^T \underline{P}_k^\perp \underline{h}_{k+1}} \underline{P}_k^\perp \underline{h}_{k+1}$$

$$= \underline{H}_k \hat{\underline{\theta}}_k + \frac{(\underline{I} - \underline{P}_k) \underline{h}_{k+1} \underline{h}_{k+1}^T \underline{P}_k^\perp \underline{x}}{\underline{h}_{k+1}^T \underline{P}_k^\perp \underline{h}_{k+1}}$$

$$= \underline{H}_k \hat{\underline{\theta}}_k + \frac{\underline{h}_{k+1} \underline{h}_{k+1}^T \underline{P}_k^\perp \underline{x}}{\underline{h}_{k+1}^T \underline{P}_k^\perp \underline{h}_{k+1}} - \frac{\underline{H}_k (\underline{H}_k^T \underline{H}_k)^{-1} \underline{H}_k^T \underline{h}_{k+1} \underline{h}_{k+1}^T \underline{P}_k^\perp \underline{x}}{\underline{h}_{k+1}^T \underline{P}_k^\perp \underline{h}_{k+1}}$$

$$= [\underline{H}_k \quad \underline{h}_{k+1}] \left[\hat{\underline{\theta}}_k - \frac{(\underline{H}_k^T \underline{H}_k)^{-1} \underline{H}_k^T \underline{h}_{k+1} \underline{h}_{k+1}^T \underline{P}_k^\perp \underline{x}}{\underline{h}_{k+1}^T \underline{P}_k^\perp \underline{h}_{k+1}} \right]$$

$$= H_{n+1} \hat{\theta}_{k+1}$$

15) Clearly, $\hat{A}_1 = \bar{x}$. From (8.28) with
 $\underline{H}_1 = \underline{1}$ $\underline{h}_2 = [1 \ r \ \dots \ r^{N-1}]^T = \underline{h}$

$$\hat{A}_2 = \hat{A}_1 - \frac{\underline{1}^T \underline{h} \underline{h}^T \underline{p}_1^+ \underline{x}}{\underline{h}^T \underline{p}_1^+ \underline{h}}$$

$$\hat{B}_2 = \frac{\underline{h}^T \underline{p}_1^+ \underline{x}}{\underline{h}^T \underline{p}_1^+ \underline{h}}$$

$$\underline{p}_1^+ = \underline{I} - \underline{1}(\underline{1}^T \underline{1})^{-1} \underline{1}^T = \underline{I} - 1/N \underline{1} \underline{1}^T$$

$$\underline{h}^T \underline{p}_1^+ \underline{x} = \underline{h}^T (\underline{x} - \bar{x} \underline{1})$$

$$= \sum_n x(n) r^n - \bar{x} \sum_n r^n$$

$$\underline{h}^T \underline{p}_1^+ \underline{h} = \underline{h}^T (\underline{I} - 1/N \underline{1} \underline{1}^T) \underline{h}$$

$$= \underline{h}^T \underline{h} - 1/N (\underline{1}^T \underline{h})^2$$

$$= \sum_n r^{2n} - 1/N (\sum_n r^n)^2$$

$$\hat{A}_2 = \bar{x} - \frac{\sum_n r^n \left[\sum_n x(n) r^n - \bar{x} \sum_n r^n \right]}{\sum_n r^{2n} - \frac{1}{N} (\sum_n r^n)^2}$$

$$\hat{B}_2 = \frac{\sum_n x(n) r^n - \bar{x} \sum_n r^n}{\sum_n r^{2n} - \frac{1}{N} (\sum_n r^n)^2}$$

$$16) \quad \underline{H}_1 = [1 \ 1 \ \dots \ 1]^T \Rightarrow \hat{A}_1 = (\underline{H}_1^T \underline{H}_1)^{-1} \underline{H}_1^T \underline{x} = \bar{x}$$

$$J_{M, \underline{H}_1} = \sum_{n=0}^{N-1} x^2(n) - N \bar{x}^2$$

$$\underline{H}_2 = \left[\begin{array}{c} \vdots \\ \vdots \\ \vdots \end{array} \right] \begin{array}{l} M \times 1 \\ (N-M) \times 1 \end{array} \quad \text{for } \underline{\theta}_2 = \begin{bmatrix} A \\ B-A \end{bmatrix}$$

$$\begin{array}{cc} \uparrow & \uparrow \\ \underline{H}_1 & \underline{h}_2 \end{array}$$

$$\underline{P}_1^\perp = \underline{I} - 1/N \underline{1} \underline{1}^T$$

$$\begin{aligned} \underline{h}_2^T \underline{P}_1^\perp \underline{x} &= \underline{h}_2^T (\underline{x} - 1/N \underline{1} \underline{1}^T \underline{x}) \\ &= \underline{h}_2^T (\underline{x} - \bar{x} \underline{1}) \\ &= \sum_{n=M}^{N-1} (x(n) - \bar{x}) = \sum_M^{N-1} x(n) - (N-M) \bar{x} \end{aligned}$$

$$\begin{aligned} \underline{h}_2^T \underline{P}_1^\perp \underline{h}_2 &= \underline{h}_2^T \underline{h}_2 - \underline{h}_2^T \underline{1} \underline{1}^T \underline{h}_2 / N \\ &= (N-M) - (N-M)^2 / N \\ &= (N-M) M / N \end{aligned}$$

$$\Rightarrow \hat{A}_2 = \hat{A}_1 - \frac{1/N \underline{1}^T \underline{h}_2 \left[\sum_M^{N-1} x(n) - (N-M) \bar{x} \right]}{(N-M) M / N}$$

$$= \bar{x} - \frac{\frac{N-M}{N} \left[\sum_M^{N-1} x(n) - (N-M) \bar{x} \right]}{(N-M) M / N}$$

$$= \bar{x} - \frac{1}{M} \sum_M^{N-1} x(n) + \frac{N}{M} \bar{x} - \bar{x} = \frac{1}{M} \sum_0^{M-1} x(n)$$

$$\hat{B-A} = \frac{\sum_M^{N-1} x(n) - (N-M) \bar{x}}{(N-M) M / N}$$

$$\text{Let } \bar{x}_1 = \frac{1}{N-M} \sum_M^{N-1} x(n)$$

$$\hat{B-A} = \frac{N}{M} (\bar{x}_1 - \bar{x})$$

$$\hat{\theta}_2 = \begin{bmatrix} \frac{1}{M} \sum_M^{N-1} x(n) \\ \frac{N}{M} (\bar{x}_1 - \bar{x}) \end{bmatrix}$$

The decrease in $J_{M,N}$ is from (8.31)

$$\begin{aligned} \frac{(\underline{h}_2^T \underline{P}_1^+ \underline{x})^2}{\underline{h}_2^T \underline{P}_1^+ \underline{h}_2} &= \underline{h}_2^T \underline{P}_1^+ \underline{x} (\hat{\theta}_2)_2 \\ &= \left(\sum_M^{N-1} x(n) - (N-M)\bar{x} \right) (\hat{\theta}_2)_2 \\ &= (N-M)(\bar{x}_1 - \bar{x}) \frac{N}{M} (\bar{x}_1 - \bar{x}) \\ &= \frac{N}{M} (N-M) (\bar{x}_1 - \bar{x})^2 \end{aligned}$$

To detect a jump (positive or negative) we test to see if the LS error decreases significantly for the second-order model or if $(\bar{x}_1 - \bar{x})^2$ is large. For no jump we expect $\bar{x}_1 \approx \bar{x}$ but for a jump or $A \neq B$, $(\bar{x}_1 - \bar{x})^2$ will be large. Also, note that the decrease in $J_{M,N}$ is just $(N-M) \frac{M}{N} (\hat{B-A})^2$.

- 17) Wish to show that $\underline{h}_i^T \underline{P}_k^+ \underline{x} = 0$
for $i = 1, 2, \dots, k$

$$\begin{bmatrix} \underline{h}_1^T \underline{P}_k^+ \underline{x} \\ \vdots \\ \underline{h}_k^T \underline{P}_k^+ \underline{x} \end{bmatrix} = \begin{bmatrix} \underline{h}_1^T \\ \vdots \\ \underline{h}_k^T \end{bmatrix} \quad \underline{P}_k^+ \underline{x} = \underline{H}_k^T \underline{P}_k^+ \underline{x}$$

$$\begin{aligned} &= \underline{H}_k^T (\underline{I} - \underline{H}_k (\underline{H}_k^T \underline{H}_k)^{-1} \underline{H}_k^T) \underline{x} \\ &= (\underline{H}_k^T - \underline{H}_k^T) \underline{x} = 0 \end{aligned}$$

$$\begin{aligned} 18) \quad J_{MIN_{k+1}} &= \underline{x}^T \underline{P}_{k+1}^+ \underline{x} \\ &= \underline{x}^T \underline{x} - \underline{x}^T \underline{P}_{k+1} \underline{x} \\ &= \underline{x}^T \underline{x} - \underline{x}^T \underline{P}_k \underline{x} - \underline{x}^T \frac{\underline{P}_k^+ \underline{h}_{k+1} \underline{h}_{k+1}^T \underline{P}_k^+ \underline{x}}{\underline{h}_{k+1}^T \underline{P}_k^+ \underline{h}_{k+1}} \\ &= J_{MIN_k} - \frac{(\underline{h}_{k+1}^T \underline{P}_k^+ \underline{x})^2}{\underline{h}_{k+1}^T \underline{P}_k^+ \underline{h}_{k+1}} \end{aligned}$$

$$\begin{aligned} 19) \quad J_{MIN}[N] &= \sum_0^N \frac{1}{\sigma_n^2} (x[n] - \hat{A}[n])^2 \\ &= \sum_0^{N-1} \frac{1}{\sigma_n^2} [x[n] - \hat{A}[n-1] - K[N](x[N] - \hat{A}[n-1])]^2 \\ &\quad + \frac{1}{\sigma_N^2} (x[N] - \hat{A}[N])^2 \\ &= \sum_0^{N-1} \frac{1}{\sigma_n^2} (x[n] - \hat{A}[n-1])^2 - 2K[N](x[N] - \hat{A}[n-1]) \\ &\quad \cdot \sum_0^{N-1} \frac{1}{\sigma_n^2} (x[n] - \hat{A}[n-1]) \\ &\quad + K^2[N](x[N] - \hat{A}[n-1])^2 \sum_0^{N-1} \frac{1}{\sigma_n^2} \\ &\quad + \frac{1}{\sigma_N^2} [x[N] - \hat{A}[n-1] - K[N](x[N] - \hat{A}[n-1])]^2 \end{aligned}$$

$$\text{But } \sum_0^{N-1} \frac{1}{\sigma_n^2} (x[n] - \hat{A}[n-1]) = 0$$

due to definition of $\hat{A}[N-1]$.

$$\begin{aligned}
 J_{MIN}[N] &= J_{MIN}[N-1] + \frac{K^2[N]}{\text{var}(\hat{A}[N-1])} (x[N] - \hat{A}[N-1])^2 \\
 &+ \frac{1}{\sigma_N^2} (x[N] - \hat{A}[N-1])^2 - 2 \frac{K[N]}{\sigma_N^2} (x[N] - \hat{A}[N-1])^2 \\
 &+ \frac{1}{\sigma_N^2} K^2[N] (x[N] - \hat{A}[N-1])^2 \\
 &= J_{MIN}[N-1] + (x[N] - \hat{A}[N-1])^2 J
 \end{aligned}$$

where $J = \frac{K^2[N]}{\text{var}(\hat{A}[N-1])} + \frac{1}{\sigma_N^2} - \frac{2K[N]}{\sigma_N^2} + \frac{K^2[N]}{\sigma_N^2}$

but from (8.38), $1 - K[N] = \frac{\sigma_N^2}{\text{var}(\hat{A}[N-1]) + \sigma_N^2}$

$$\begin{aligned}
 J &= \frac{\text{var}(\hat{A}[N-1])}{(\text{var}(\hat{A}[N-1]) + \sigma_N^2)^2} + \frac{1}{\sigma_N^2} \frac{\sigma_N^4}{(\text{var}(\hat{A}[N-1]) + \sigma_N^2)^2} \\
 &= \frac{1}{\text{var}(\hat{A}[N-1]) + \sigma_N^2}
 \end{aligned}$$

20) $H[N] = \begin{bmatrix} 1 \\ r \\ \vdots \\ r^n \end{bmatrix} \Rightarrow h[n] = r^n$

$$\hat{A}[n] = \hat{A}[n-1] + K[n] (x[n] - r^n \hat{A}[n-1])$$

from (8.46)

Now $\sigma_n^2 = \sigma^2 = 1$ and $\Sigma(n) = \text{var}(\hat{A}(n))$

$$\Rightarrow K(n) = \frac{\text{var}(\hat{A}(n-1)) r^n}{1 + r^{2n} \text{var}(\hat{A}(n-1))} \quad \text{from (8.47)}$$

$$\text{var}(\hat{A}(n)) = (1 - K(n) r^n) \text{var}(\hat{A}(n-1)) \quad \text{from (8.48)}$$

To find the variance explicitly let

$$\nu_n = \text{var}(\hat{A}(n))$$

$$\begin{aligned} \nu_n &= \left[1 - \frac{\nu_{n-1} r^{2n}}{1 + r^{2n} \nu_{n-1}} \right] \nu_{n-1} \\ &= \frac{\nu_{n-1}}{1 + r^{2n} \nu_{n-1}} \end{aligned}$$

Now let $\nu_0 = 1$ as was given

$$\begin{aligned} \nu_1 &= \frac{1}{1+r^2} \quad \nu_2 = \frac{\frac{1}{1+r^2}}{1+r^4\left(\frac{1}{1+r^2}\right)} \\ &= \frac{1}{1+r^2+r^4} \end{aligned}$$

or in general

$$\nu_n = \text{var}(\hat{A}(n)) = \frac{1}{\sum_{k=0}^n r^{2k}}$$

$$21) \quad K(N) = \frac{\text{var}(\hat{A}(N-1))}{\text{var}(\hat{A}(N-1)) + \sigma_N^2}$$

$$\text{var}(\hat{A}[N]) = (1 - K[N]) \text{var}(\hat{A}[N-1])$$

$$\text{Let } N_N = \text{var}(\hat{A}[N])$$

$$N_N = \left(1 - \frac{N_{N-1}}{N_{N-1} + \sigma_N^2}\right) N_{N-1}$$

$$= \frac{\sigma_N^2 N_{N-1}}{N_{N-1} + \sigma_N^2}$$

$$\frac{1}{N_N} = \frac{N_{N-1} + \sigma_N^2}{\sigma_N^2 N_{N-1}} = \frac{1}{\sigma_N^2} + \frac{1}{N_{N-1}}$$

$$= \frac{1}{r^N} + \frac{1}{N_{N-1}}$$

Since $\frac{1}{N_0} = 1$, we have

$$\frac{1}{N_N} = \sum_{n=0}^N 1/r^n \quad \text{or} \quad N_N = \frac{1}{\sum_{n=0}^N 1/r^n}$$

$$\begin{aligned} K[N] &= \frac{1 / \sum_{n=0}^N 1/r^n}{1 / \sum_{n=0}^N 1/r^n + r^N} = \frac{1}{1 + r^N \sum_{n=0}^N 1/r^n} \\ &= \frac{1}{1 + \sum_{n=0}^N r^{N-n}} = \frac{1}{1 + \sum_{n=0}^N r^n} \end{aligned}$$

If $r = 1$, $\text{var}(\hat{A}[N]) \rightarrow 0$ as $N \rightarrow \infty$

$K[N] \rightarrow 0$ as $N \rightarrow \infty$

If $0 < r < 1$, $\text{var}(\hat{A}[N]) \rightarrow 0$ as $N \rightarrow \infty$
 $K[N] \rightarrow \text{Constant}$ as $N \rightarrow \infty$

If $r > 1$, $\text{var}(\hat{A}[N]) \rightarrow \frac{r-1}{r}$ as $N \rightarrow \infty$

since $\sum_{n=0}^{\infty} 1/r^n = \frac{1}{1-1/r} = \frac{r}{r-1}$

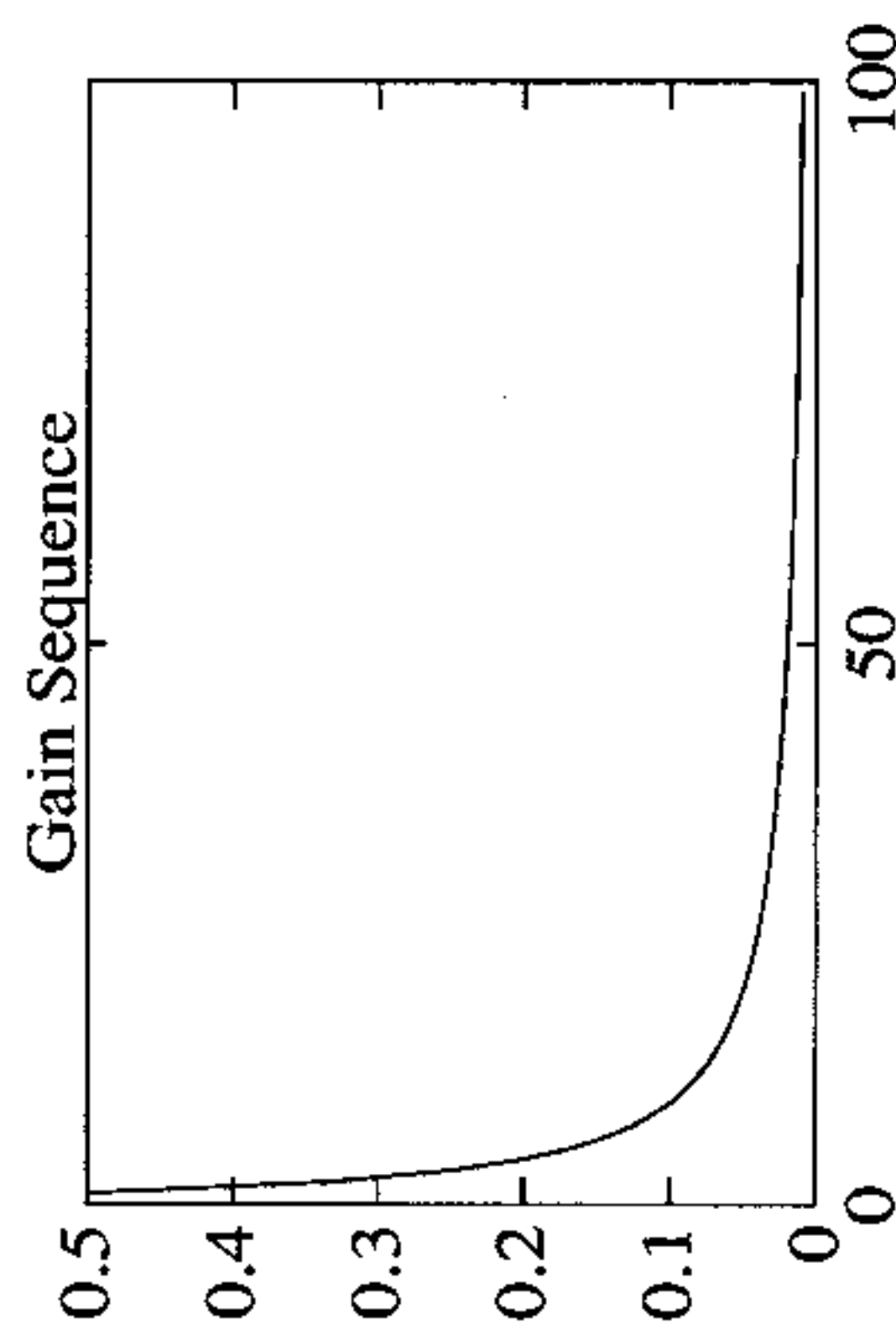
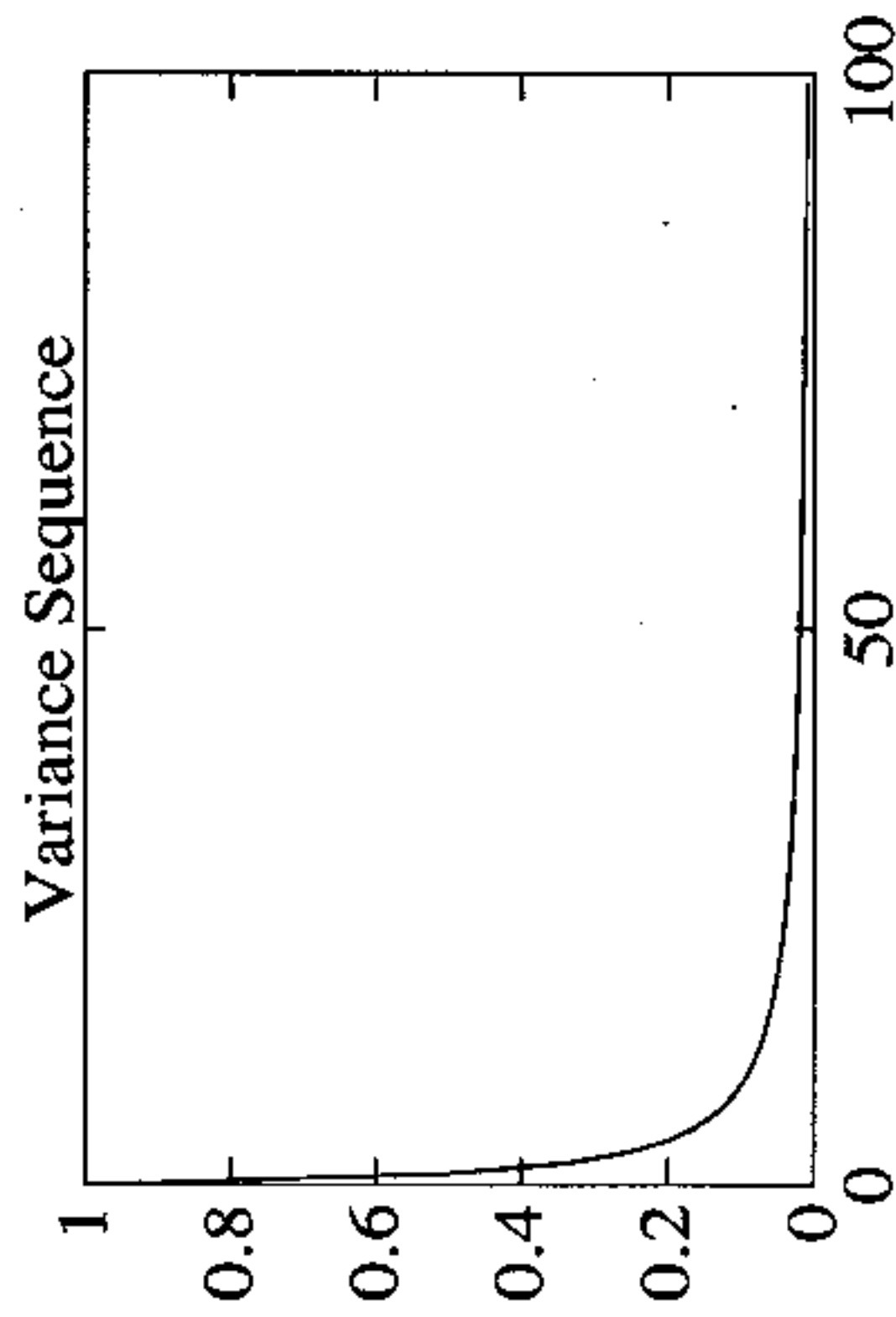
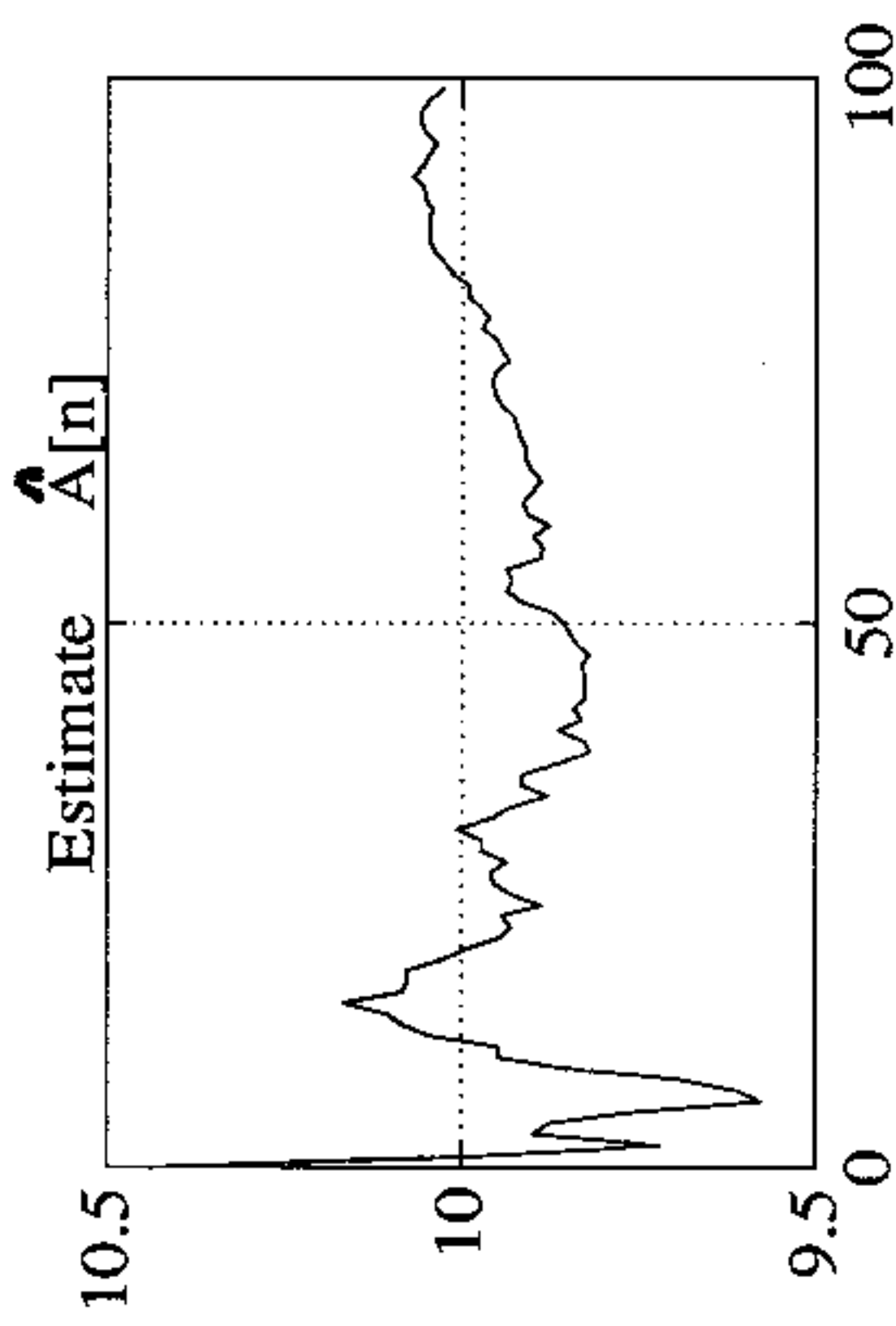
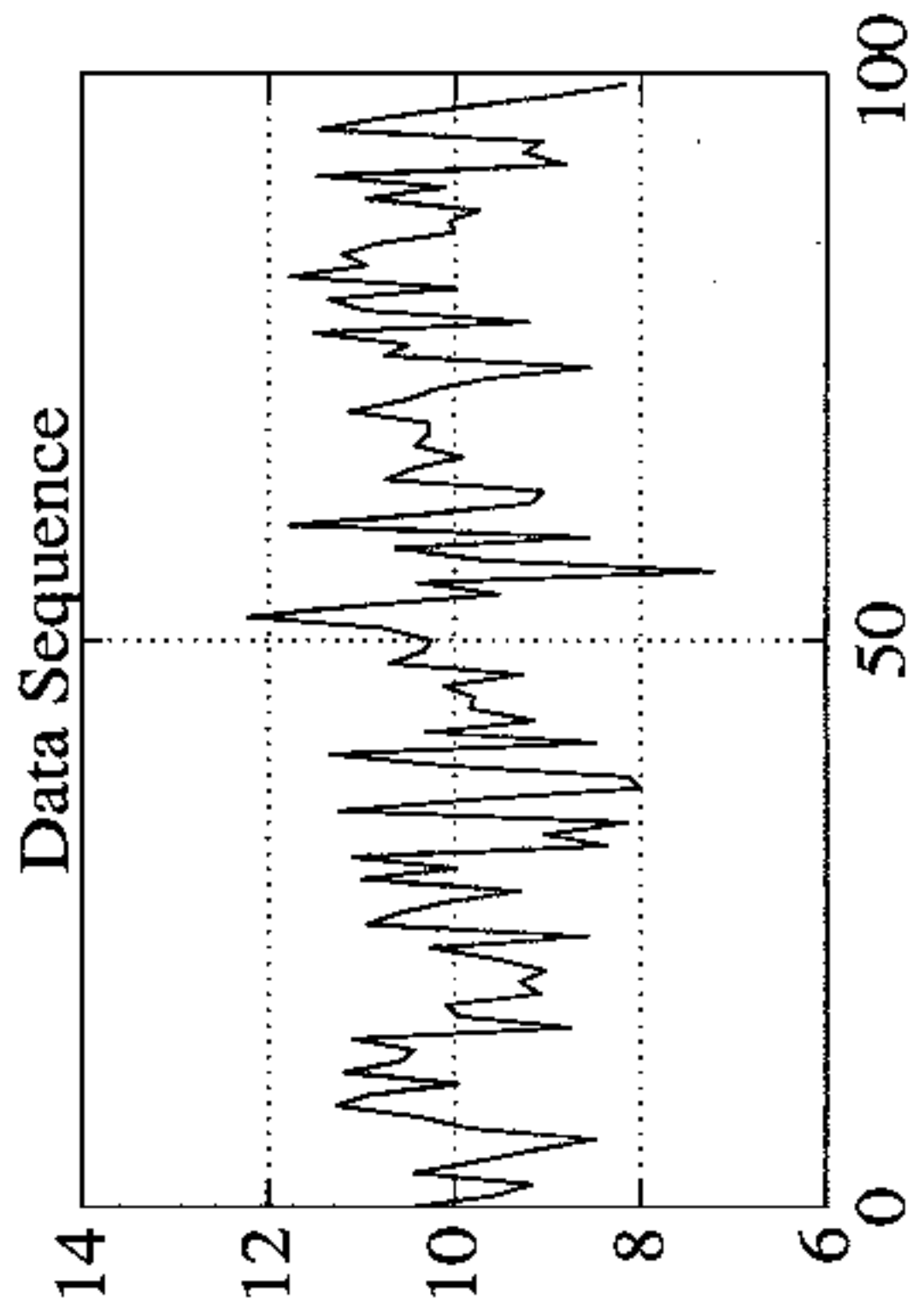
and $K[N] \rightarrow 0$ as $N \rightarrow \infty$. In this case the data become so noisy that the gain $\rightarrow 0$ (we do not use the data) and hence the variance does not go to zero.

22) See plots on next few pages. Compare these results to those of the previous problem.

$$23) \hat{\theta}_r(n) = \left(\sum_{k=-p}^{-1} \frac{1}{\sigma_k^2} \underline{h}(k) \underline{h}^T(k) + \sum_{k=0}^n \frac{1}{\sigma_k^2} \underline{h}(k) \underline{h}^T(k) \right)^{-1} \\ \cdot \left(\sum_{k=-p}^{-1} \frac{1}{\sigma_k^2} x(k) \underline{h}^T(k) + \sum_{k=0}^n \frac{1}{\sigma_k^2} x(k) \underline{h}^T(k) \right)$$

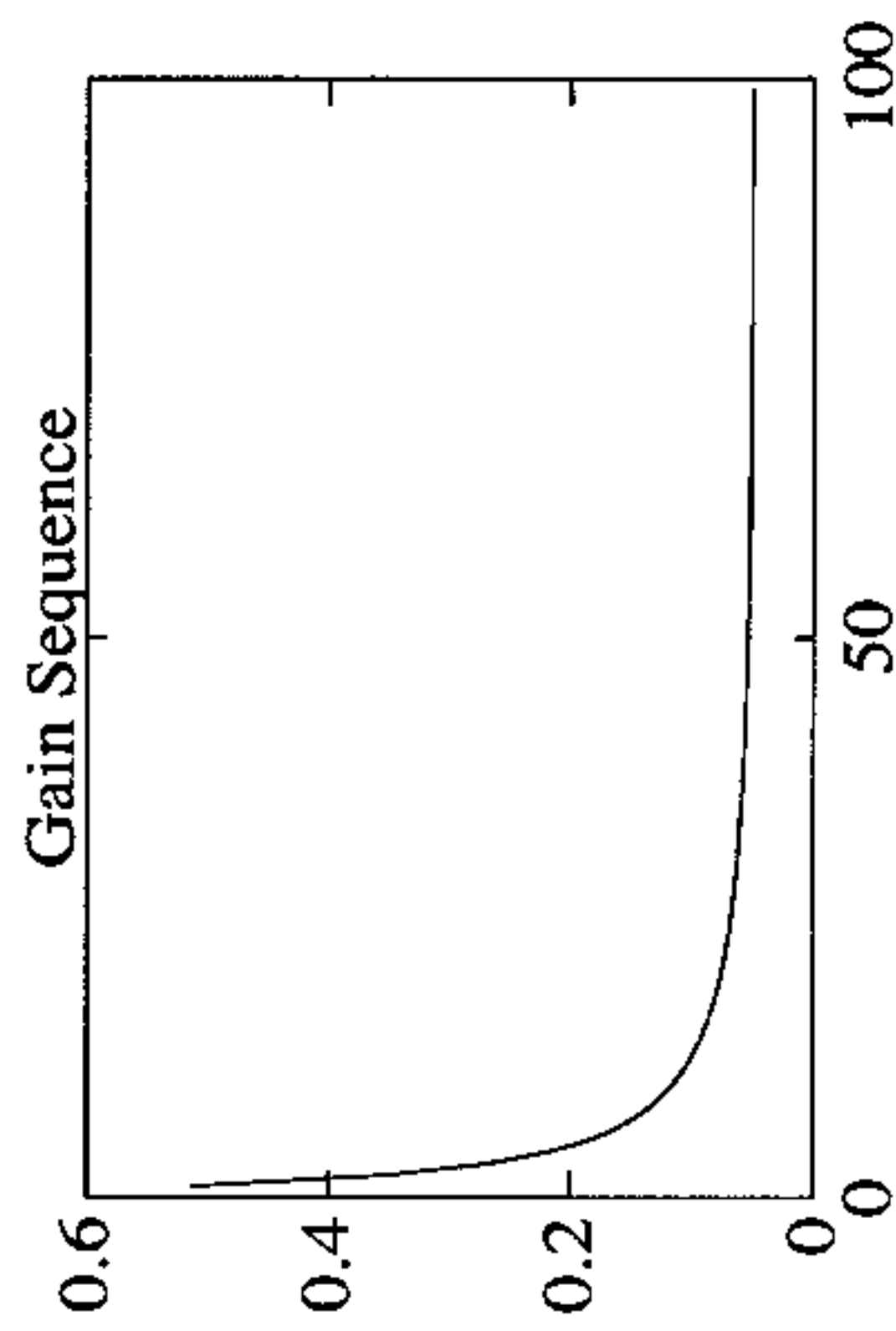
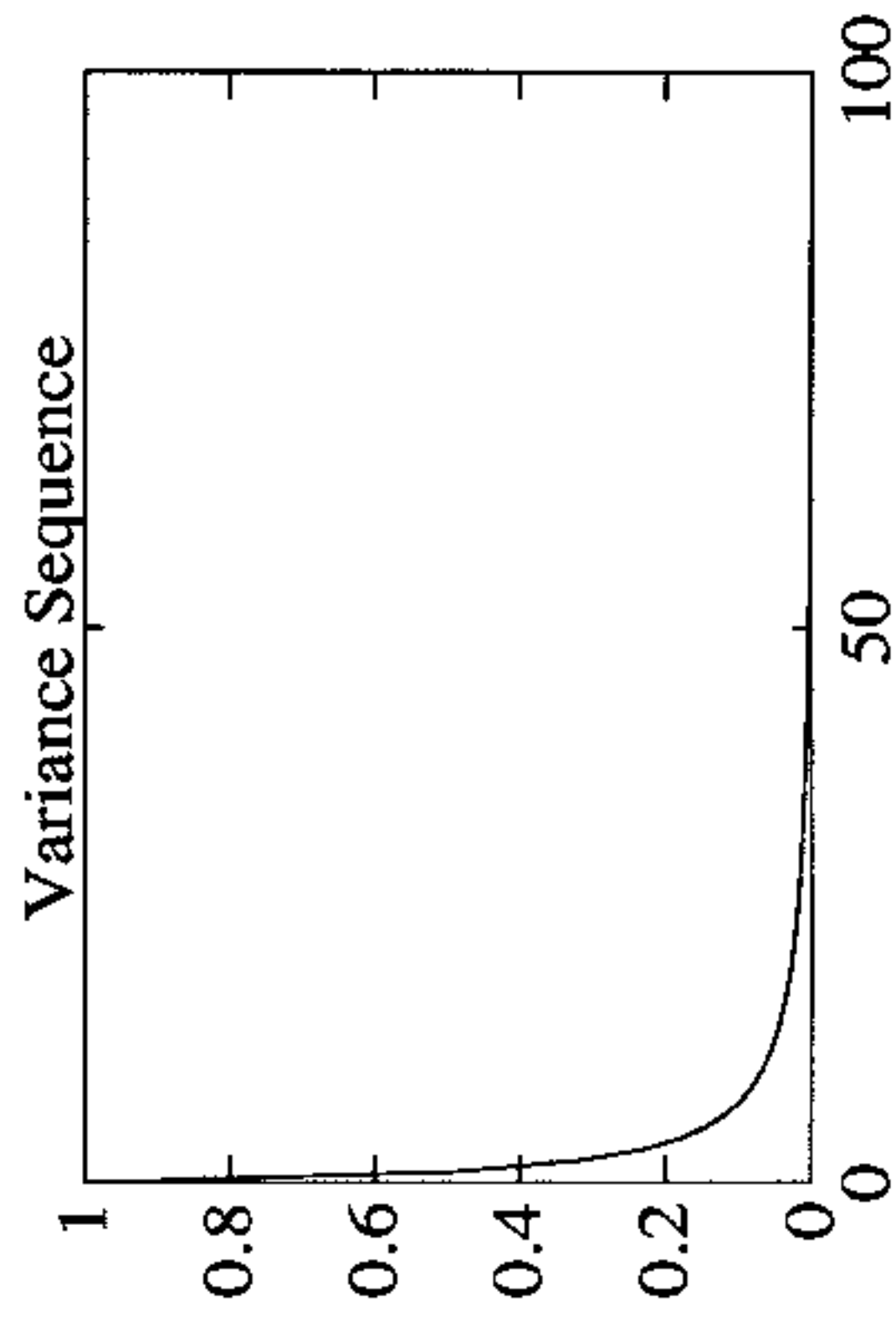
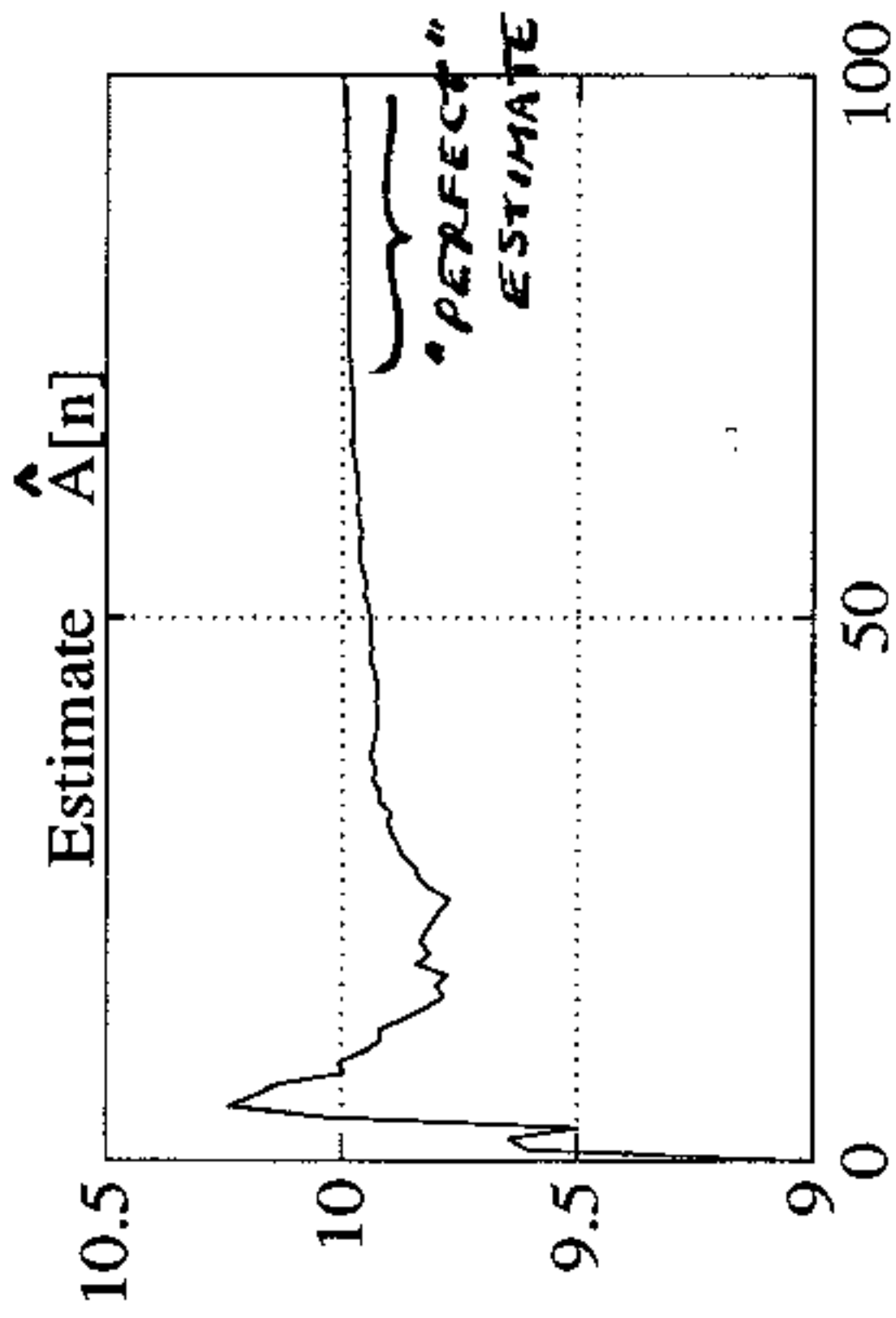
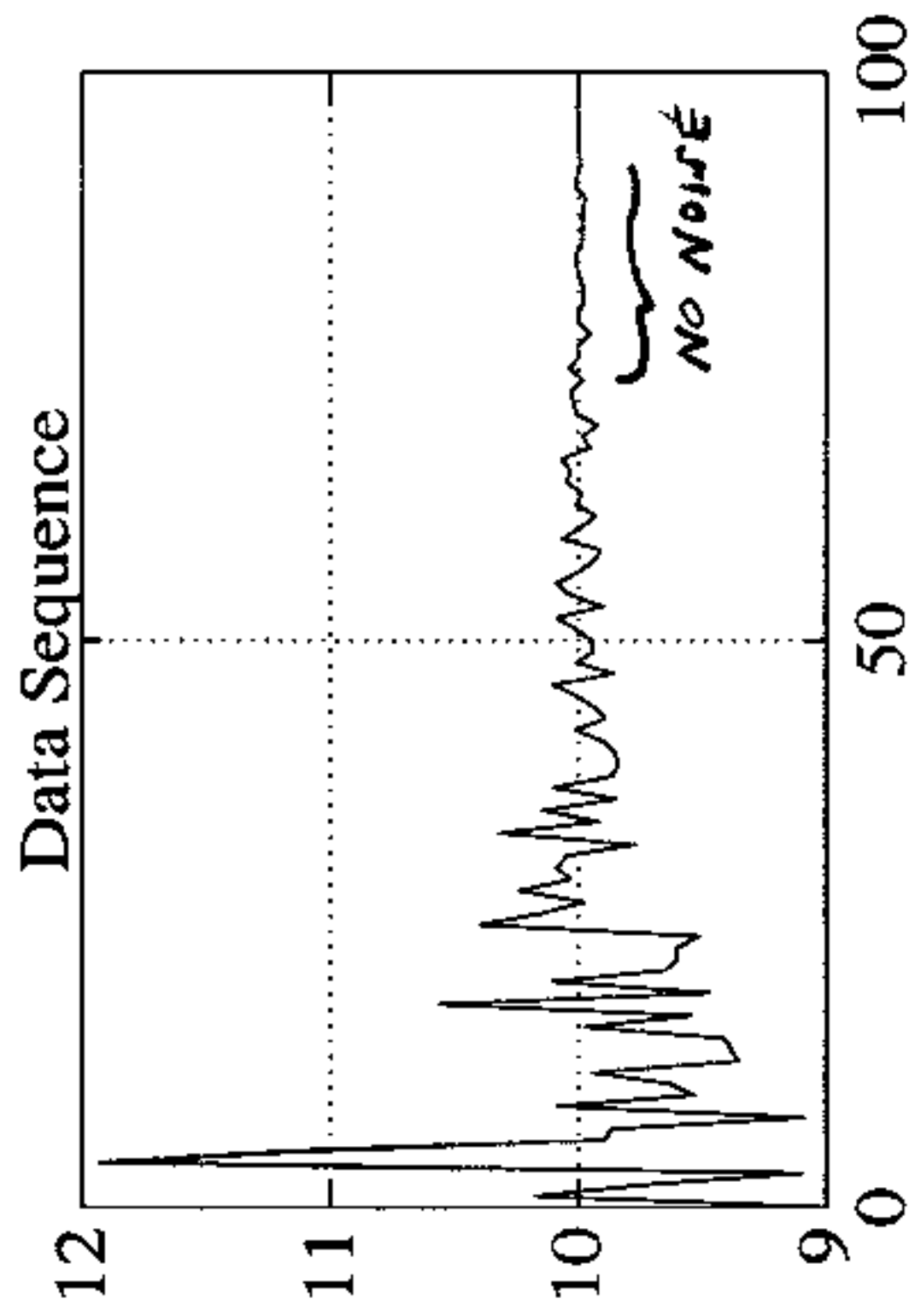
$$\text{But } \hat{\theta}(-1) = \left(\sum_{k=-p}^{-1} \frac{1}{\sigma_k^2} \underline{h}(k) \underline{h}^T(k) \right)^{-1} \\ \cdot \left(\sum_{k=-p}^{-1} \frac{1}{\sigma_k^2} x(k) \underline{h}^T(k) \right)$$

$$\underline{\Sigma}(-1) = \left(\sum_{k=-p}^{-1} \frac{1}{\sigma_k^2} \underline{h}(k) \underline{h}^T(k) \right)^{-1}$$



$$\underline{r = 1.00}$$

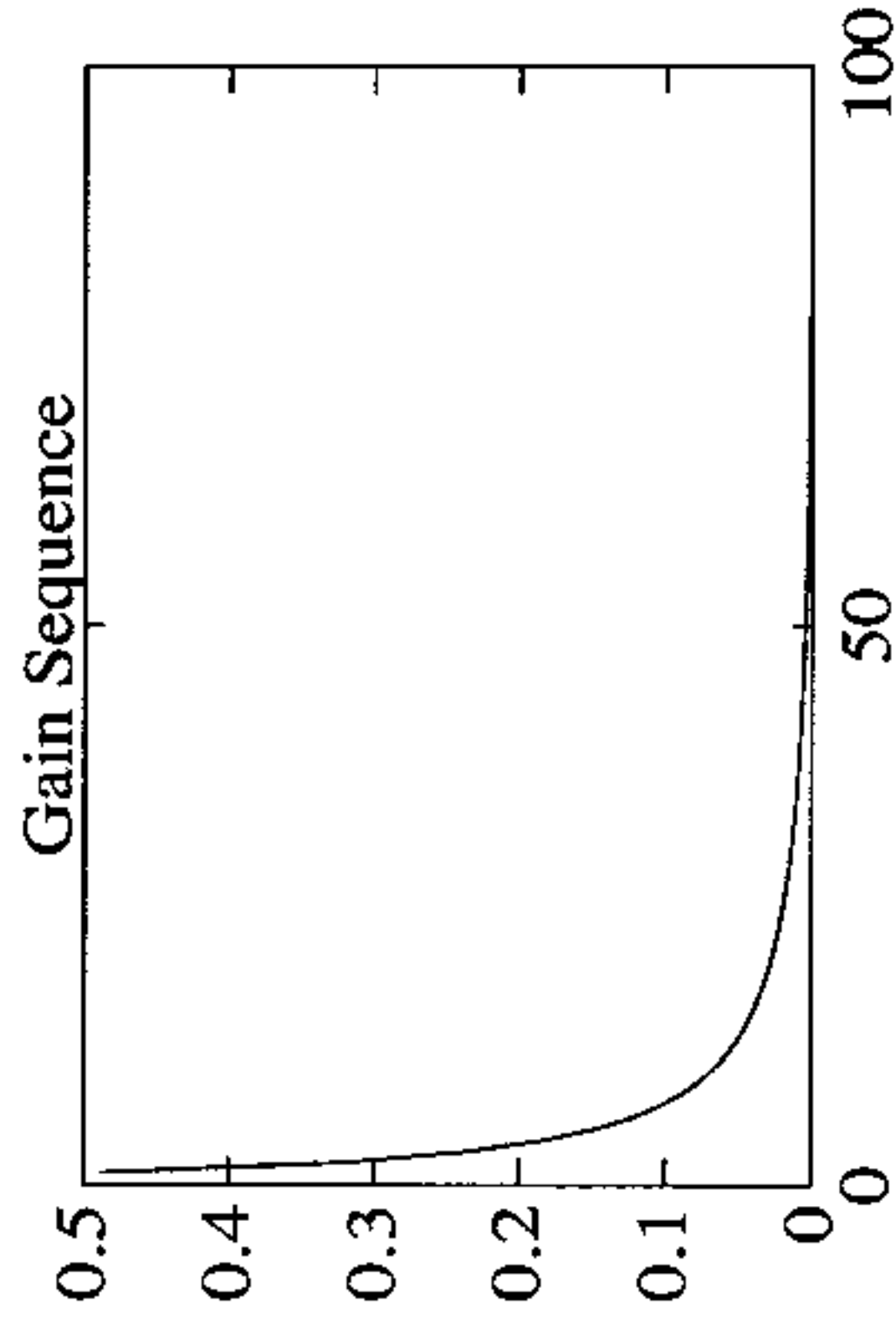
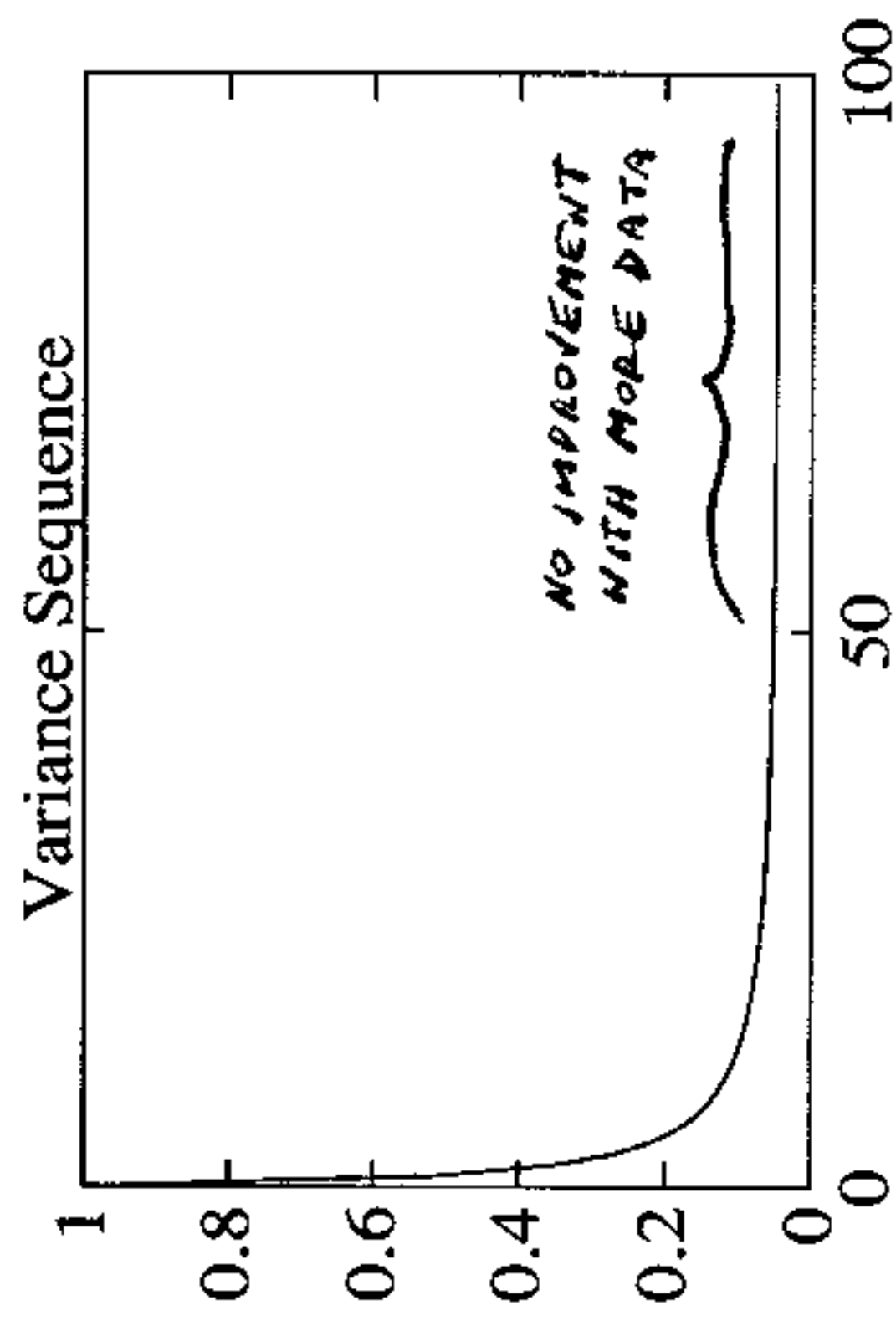
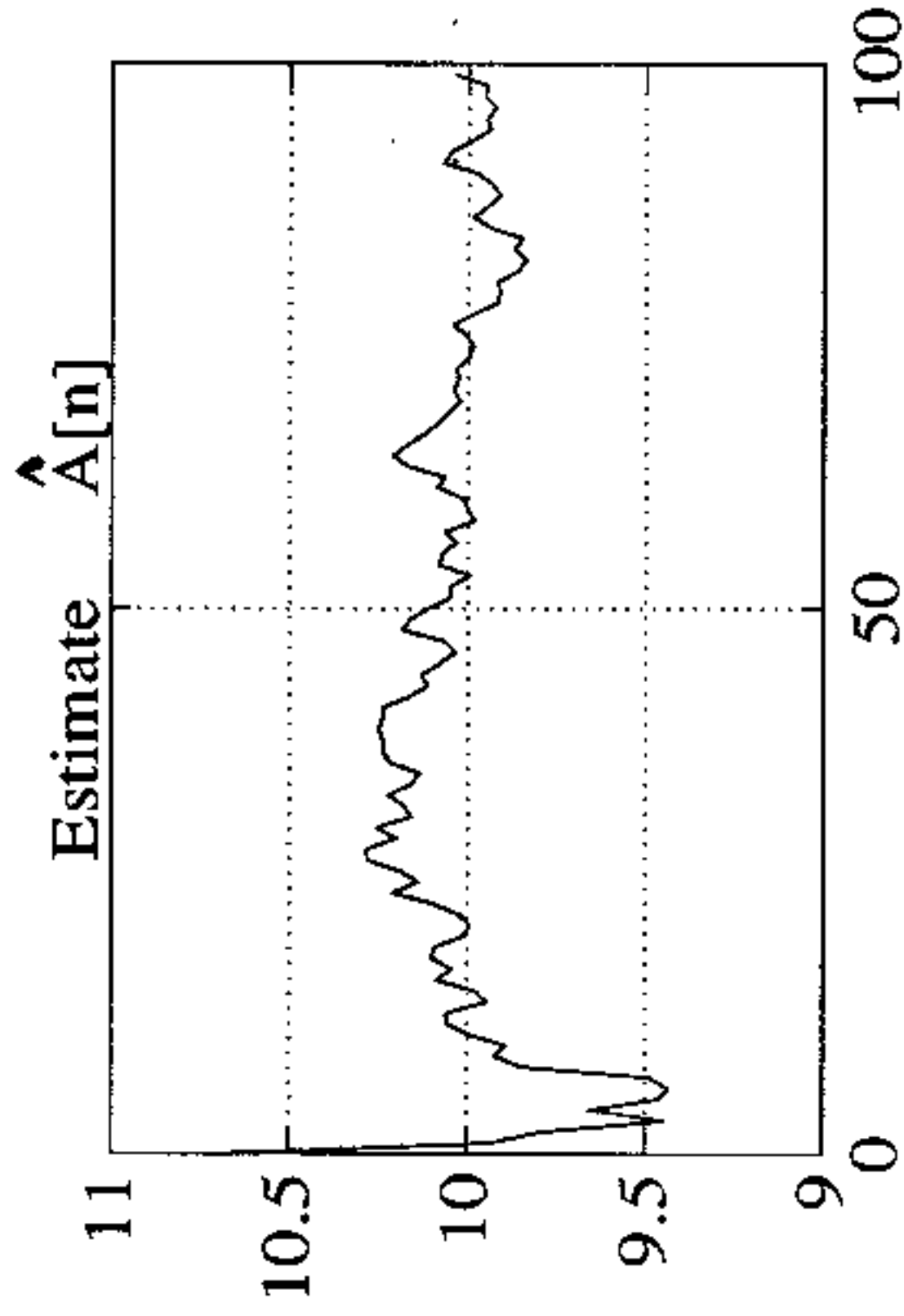
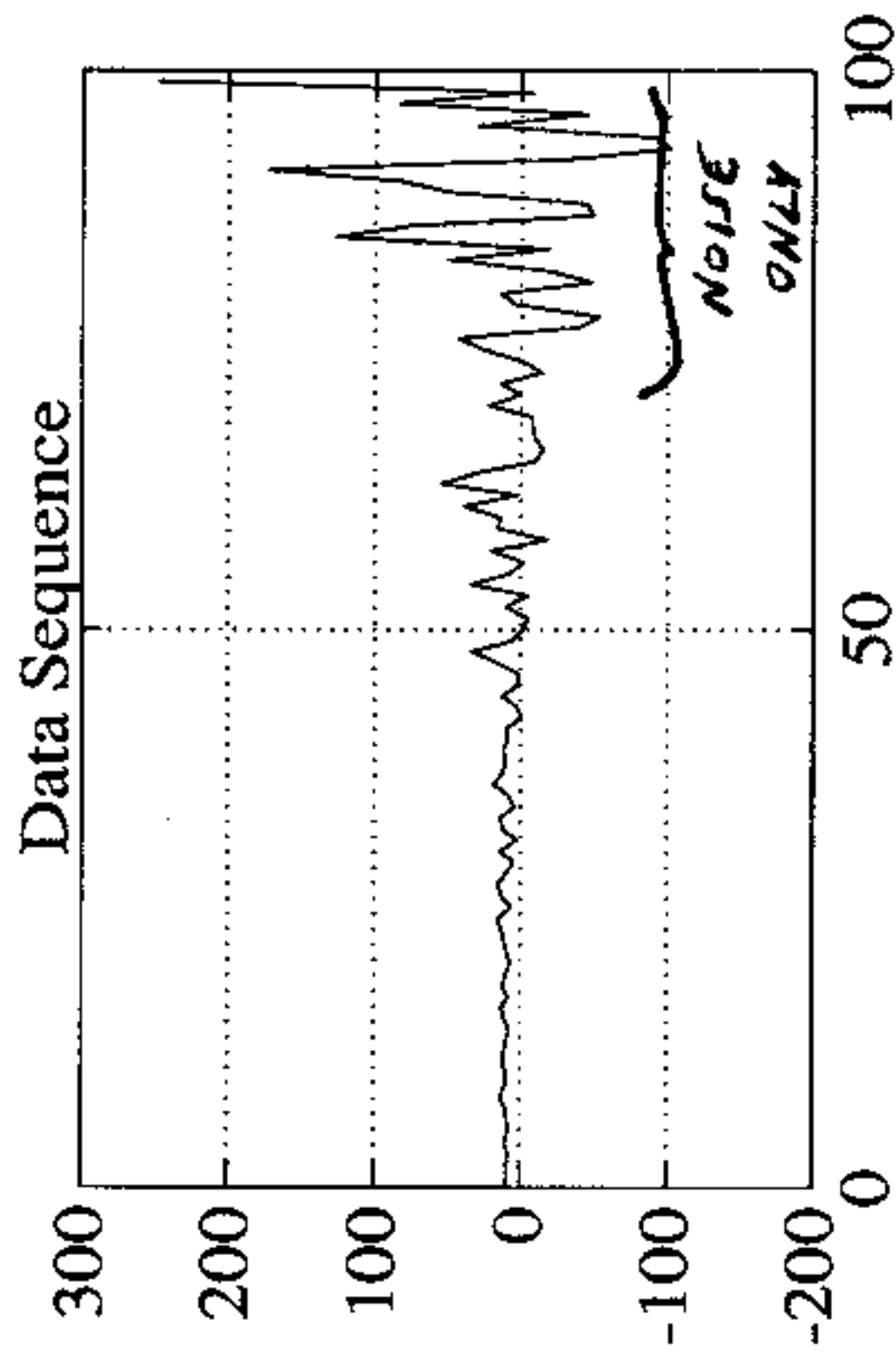
ob. 8.22



$r = 0.95$

120

Prob. 8.22



$r = 1.05$

$$\Rightarrow \hat{\underline{\theta}}_S(n) = \left(\underline{\Sigma}^{-1}[-1] + \sum_{k=0}^n \frac{1}{\sigma_k^2} \underline{h}(k) \underline{h}^T(k) \right)^{-1} \\ \cdot \left(\underline{\Sigma}^{-1}[-1] \hat{\underline{\theta}}[-1] + \sum_{k=0}^n \frac{1}{\sigma_k^2} x(k) \underline{h}^T(k) \right)$$

Now for $n \geq p$, $\sum_{k=0}^n \frac{1}{\sigma_k^2} \underline{h}(k) \underline{h}^T(k)$ will

be invertible so we can let $\alpha \rightarrow \infty$. Then,
 $\underline{\Sigma}^{-1}[-1] \rightarrow 0$ to yield

$$\hat{\underline{\theta}}_S(n) = \left(\sum_{k=0}^n \frac{1}{\sigma_k^2} \underline{h}(k) \underline{h}^T(k) \right)^{-1} \\ \cdot \left(\sum_{k=0}^n \frac{1}{\sigma_k^2} x(k) \underline{h}^T(k) \right) \\ = \hat{\underline{\theta}}_B(n)$$

Note that if $\underline{\Sigma}[-1] \rightarrow \infty$, the choice of $\hat{\underline{\theta}}[-1]$ is immaterial.

24) Projecting $\underline{x}$ onto subspace spanned by $\underline{h}_1$ and $\underline{h}_2$ produces $\hat{\underline{z}} = \underline{H}(\underline{H}^T \underline{H})^{-1} \underline{H}^T \underline{x}$ where $\underline{H} = [\underline{h}_1, \underline{h}_2]$. Now the constraint subspace is spanned by $[1, 1, 0]^T$. The projection onto this subspace is

$$\hat{\underline{z}}_c = \underline{e}_c (\underline{e}_c^T \underline{e}_c)^{-1} \underline{e}_c^T \hat{\underline{z}} \quad \underline{e}_c = \begin{bmatrix} 1 \\ 1 \\ 0 \end{bmatrix} \\ = \frac{1}{2} \begin{bmatrix} 1 \\ 1 \\ 0 \end{bmatrix} (1, 1, 0) \begin{bmatrix} x(0) \\ x(1) \\ 0 \end{bmatrix}$$

$$= \frac{1}{2} \begin{bmatrix} 1 & 1 & 0 \\ 1 & 1 & 0 \\ 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} x_{10} \\ x_{11} \\ 0 \end{bmatrix} = \begin{bmatrix} \frac{1}{2}(x_{10} + x_{11}) \\ \frac{1}{2}(x_{10} + x_{11}) \\ 0 \end{bmatrix}$$

25) $H = \begin{bmatrix} 1 & 1 \\ 1 & -1 \\ \vdots & \vdots \\ 1 & -1 \end{bmatrix} \Rightarrow$ columns of H are orthogonal

$$\begin{aligned} \hat{\theta} &= (H^T H)^{-1} H^T x = \begin{bmatrix} N & 0 \\ 0 & N \end{bmatrix}^{-1} \begin{bmatrix} 1 & 1 \\ 1 & -1 \\ \vdots & \vdots \\ 1 & -1 \end{bmatrix}^T x \\ &= \begin{bmatrix} \frac{1}{N} \sum_{n=0}^{N-1} x(n) \\ \frac{1}{N} \sum_{n=0}^{N-1} (-1)^n x(n) \end{bmatrix} \end{aligned}$$

Now, if $A = B$ we have $(1 - 1)\theta = 0$

or $A = [1 \ -1]$, $b = 0$ and thus from (P. 52)

$$\begin{aligned} \hat{\theta}_c &= \hat{\theta} - (H^T H)^{-1} A^T (A (H^T H)^{-1} A^T)^{-1} A \hat{\theta} \\ &= \hat{\theta} - \frac{1}{N} A^T \left(\frac{1}{N} A A^T \right)^{-1} A \hat{\theta} \\ &= (\mathbf{I} - A^T (A A^T)^{-1} A) \hat{\theta} \end{aligned}$$

$$\begin{aligned} A^T (A A^T)^{-1} A &= \begin{bmatrix} 1 & -1 \end{bmatrix} \left(\begin{bmatrix} 1 & -1 \end{bmatrix} \begin{bmatrix} 1 \\ -1 \end{bmatrix} \right)^{-1} \begin{bmatrix} 1 & -1 \end{bmatrix} \\ &= \frac{1}{2} \begin{bmatrix} 1 & -1 \\ -1 & 1 \end{bmatrix} \end{aligned}$$

$$\mathbf{I} - A^T (A A^T)^{-1} A = \frac{1}{2} \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix}$$

$$\hat{\theta}_c = \begin{bmatrix} \frac{1}{2} \left(\bar{x} + \frac{1}{N} \sum_{n=0}^{N-1} (-1)^n x(n) \right) \\ \frac{1}{2} \left(\bar{x} + \frac{1}{N} \sum_{n=0}^{N-1} (-1)^n x(n) \right) \end{bmatrix}$$

$$\text{or } \hat{A}_c = \hat{B}_c = \frac{1}{2N} \sum_{\substack{n=0 \\ n \text{ even}}}^{N-1} x[n]$$

Makes sense, since if $A=B$

$$S[n] = \begin{cases} A & n \text{ even} \\ 0 & n \text{ odd} \end{cases}$$

26) By assumption

$$h(g^{-1}(\hat{\alpha})) \leq h(g^{-1}(\alpha)) \text{ for all } \alpha$$

$$\Rightarrow h(\theta_0) \leq h(\theta) \text{ for all } \theta$$

where $\theta_0 = g^{-1}(\hat{\alpha})$

But then θ_0 minimizes $h(\theta)$, i.e., $\theta_0 = \hat{\theta}$

$$\Rightarrow \hat{\theta} = g^{-1}(\hat{\alpha})$$

27) From (8.61)

$$\theta_{k+1} = \theta_k + \left(\underline{H}^T(\theta_k) \underline{H}(\theta_k) - \sum_{n=0}^{N-1} G_n(\theta_k) (x[n] - e^{\theta_k}) \right)^{-1} \cdot \underline{H}^T(\theta_k) (\underline{x} - e^{\theta_k})$$

$$[\underline{H}(\theta)]_i = \frac{\partial S[i]}{\partial \theta} = e^{\theta} \quad i = 0, 1, \dots, N-1$$

$$[G_n(\theta)]_{ii} = \frac{\partial^2 S[n]}{\partial \theta^2} = e^{\theta}$$

$$\Rightarrow \underline{H}(\theta) = e^{\theta} \underline{1} \quad G_n(\theta) = e^{\theta}$$

$$\theta_{k+1} = \theta_k + \left(N e^{2\theta_k} - \sum_{n=0}^{N-1} e^{\theta_k} (x[n] - e^{\theta_k}) \right)^{-1} \\ \cdot e^{\theta_k} \mathbf{1}^T (\mathbf{x} - e^{\theta_k} \mathbf{1})$$

$$= \theta_k + \frac{e^{\theta_k} (N \bar{x} - N e^{\theta_k})}{N e^{2\theta_k} - N \bar{x} e^{\theta_k} + N e^{2\theta_k}}$$

$$= \theta_k + \frac{\bar{x} - e^{\theta_k}}{2 e^{\theta_k} - \bar{x}}$$

To find analytically, let $\alpha = e^{\theta} \Rightarrow \hat{\alpha} = \bar{x}$
and from Prob 8.26 $\hat{\theta} = \ln \bar{x}$.

$$28) \text{ Since } H(z) = \frac{B(z)}{A(z)} \\ = B(z) \frac{1}{A(z)}$$

$$h[n] = b[n] * g[n] \\ = \sum_{k=0}^n b[k] g[n-k]$$

Since $g[n]$ is causal. In matrix form
for $n = 0, 1, \dots, N-1$ we have $\underline{s} = \underline{G} \underline{b}$.

$$\mathcal{J}(\underline{a}, \underline{b}) = (\underline{x} - \underline{s})^T (\underline{x} - \underline{s}) \\ = (\underline{x} - \underline{G} \underline{b})^T (\underline{x} - \underline{G} \underline{b}) \\ \Rightarrow \hat{\underline{b}} = (\underline{G}^T \underline{G})^{-1} \underline{G}^T \underline{x}$$

and $J(\underline{a}, \hat{\underline{b}}) = \underline{x}^T (\underline{I} - \underline{G}(\underline{G}^T \underline{G})^{-1} \underline{G}^T) \underline{x}$
from (P. 11)

Now $G(z)/A(z) = 1 \Rightarrow g[n] \star a[n] = \delta[n]$

$$\begin{aligned} [A^T G]_{ij} &= \sum_{k=1}^N [A^T]_{ik} [G]_{kj} \\ &= \sum_{k=1}^N a[p+i-k] g[k-j] \\ &= \sum_{k=-\infty}^{\infty} a[p+i-k] g[k-j] \end{aligned}$$

$i = 1, 2, \dots, N-p$
 $j = 1, 2, \dots, p+1$
 $= 1, 2, \dots, p$

where $a[n] = 0$ for $n < 0, n > p$
 $g[n] = 0$ for $n < 0$

$$= \sum_{l=-\infty}^{\infty} g[l] a[p+i-j-l]$$

$$= g[n] \star a[n] \big|_{n=p+i-j}$$

$$= \delta[p+i-j]$$

Since $1 \leq i \leq N-p, 1 \leq j \leq p,$
 $p+i-j > 0$ for all i, j

$$\Rightarrow A^T G = 0$$

$$\underline{L} (\underline{L}^T \underline{L})^{-1} \underline{L}^T = \underline{I}$$

$$[\underline{A} \ \underline{G}] \left(\begin{bmatrix} \underline{A}^T \\ \underline{G}^T \end{bmatrix} [\underline{A} \ \underline{G}] \right)^{-1} \begin{bmatrix} \underline{A}^T \\ \underline{G}^T \end{bmatrix} = \underline{I}$$

$$[\underline{A} \ \underline{G}] \begin{bmatrix} \underline{A}^T \underline{A} & 0 \\ 0 & \underline{G}^T \underline{G} \end{bmatrix}^{-1} \begin{bmatrix} \underline{A}^T \\ \underline{G}^T \end{bmatrix} = \underline{I}$$

$$\Rightarrow \underline{A} (\underline{A}^T \underline{A})^{-1} \underline{A}^T + \underline{G} (\underline{G}^T \underline{G})^{-1} \underline{G}^T = \underline{I}$$

29) If $f_{0,k+1} = f_{0,k}$, $\phi_{k+1} = \phi_k$, we have

$$\sum_{n=-M}^M n x[n] \sin(2\pi \hat{f}_0 n + \hat{\phi}) = 0$$

$$\sum_{n=-M}^M x[n] \sin(2\pi \hat{f}_0 n + \hat{\phi}) = 0$$

At high SNR we have

$$\sum_n n \cos(2\pi f_0 n + \phi) \sin(2\pi \hat{f}_0 n + \hat{\phi}) = 0$$

$$\sum_n \cos(2\pi f_0 n + \phi) \sin(2\pi \hat{f}_0 n + \hat{\phi}) = 0$$

$$\frac{1}{2} \sum_n n \left[\sin(2\pi (f_0 + \hat{f}_0) n + \phi + \hat{\phi}) + \sin(2\pi (\hat{f}_0 - f_0) n + (\hat{\phi} - \phi)) \right] = 0$$

$$\frac{1}{2} \sum_n \left[\sin(2\pi (f_0 + \hat{f}_0) n + \phi + \hat{\phi}) + \sin(2\pi (\hat{f}_0 - f_0) n + (\hat{\phi} - \phi)) \right] = 0$$

Neglecting the high frequency term we have

$$\sum_n n \sin(2\pi(\hat{f}_0 - f_0)n + \hat{\phi} - \phi) = 0$$

$$\sum_n \sin(2\pi(\hat{f}_0 - f_0)n + \hat{\phi} - \phi) = 0$$

for which the solution is $\hat{f}_0 = f_0, \hat{\phi} = \phi$.

Chapter 9

$$1) E(X) = \int_0^{\infty} \frac{x^2}{\sigma^2} e^{-\frac{1}{2} x^2 / \sigma^2} dx = \sqrt{\pi/2} \sigma$$

$$\Rightarrow \sigma^2 = \frac{2}{\pi} E^2(X)$$

$$\hat{\sigma}^2 = \frac{2}{\pi} \left(\frac{1}{N} \sum_{n=0}^{N-1} X(n) \right)^2 = \frac{2}{\pi} \bar{X}^2$$

2) $E(X) = 0$ since the PDF is even.
Try $E(X^2)$

$$\begin{aligned} E(X^2) &= \int_{-\infty}^{\infty} x^2 \frac{1}{\sqrt{2}\sigma} e^{-\sqrt{2}|x|/\sigma} dx \\ &= \frac{2}{\sqrt{2}\sigma} \int_0^{\infty} x^2 e^{-\sqrt{2}x/\sigma} dx = \sigma^2 \end{aligned}$$

$$\Rightarrow \sigma = \sqrt{E(X^2)} \quad \text{and} \quad \hat{\sigma} = \sqrt{\frac{1}{N} \sum_{n=0}^{N-1} X^2(n)}$$

3) $\rho = \cos(u, v)$ where $\underline{x} = \begin{bmatrix} u \\ v \end{bmatrix}$
 $= E(uv)$

$$\Rightarrow \hat{\rho} = \frac{1}{N} \sum_{n=0}^{N-1} u_n v_n \quad \text{where} \quad \underline{x}(n) = \begin{bmatrix} u_n \\ v_n \end{bmatrix}$$

The cubic equation was found in Prob. 7.11.
Clearly, the method of moments estimator is much simpler to find and implement.

4) Since $\mu = E(x)$, $\sigma^2 = \text{var}(x)$

$$\hat{\mu} = \frac{1}{N} \sum_{n=0}^{N-1} x[n] \quad \hat{\sigma}^2 = \frac{1}{N} \sum_{n=0}^{N-1} (x[n] - \hat{\mu})^2$$

5) Replace the theoretical ACF by $\hat{r}_{xx}[n]$ where

$$\hat{r}_{xx}[n] = \frac{1}{N} \sum_{n=0}^{N-1-|n|} x[n] x[n+|n|]$$

$$\Rightarrow \hat{r}_{xx}[n] = - \sum_{k=1}^p a[k] \hat{r}_{xx}[n-k] \quad n > q$$

Choosing $n = q+1, \dots, q+p$ yields the linear equations

$$\begin{bmatrix} \hat{r}_{xx}[q] & \dots & \hat{r}_{xx}[q-p+1] \\ \hat{r}_{xx}[q+1] & \dots & \hat{r}_{xx}[q-p+2] \\ \vdots & & \vdots \\ \hat{r}_{xx}[q+p-1] & \dots & \hat{r}_{xx}[q] \end{bmatrix} \begin{bmatrix} a[1] \\ \vdots \\ a[p] \end{bmatrix} = - \begin{bmatrix} \hat{r}_{xx}[q+1] \\ \vdots \\ \hat{r}_{xx}[q+p] \end{bmatrix}$$

which can be solved for the $a[k]$'s.

6) Let $\theta = A^2$ so that $\hat{\theta} = g(\bar{x})$ where $g(x) = x^2$.

$T = \bar{x}$ so from (9.15)

$$E(\hat{\theta}) = g(E(T)) = g(A) = A^2$$

From (9.16)

$$\text{var}(\hat{\theta}) = \left(\frac{\partial g}{\partial T} \bigg|_{T=A} \right)^2 \text{var}(T)$$

$$= (2T|_{T=A})^2 \text{var}(T)$$

$$= (2A)^2 \sigma^2/N = 4A^2 \sigma^2/N$$

$$7) \quad E(X(n)) = \cos \phi$$

$$\Rightarrow \phi = \arccos E(X(n))$$

or

$$\hat{\phi} = \arccos \left(\frac{1}{N} \sum_{n=0}^{N-1} X(n) \right)$$

$$\hat{\phi} = h(\underline{w}) \quad \text{where}$$

$$h(\underline{w}) = \arccos \left[\frac{1}{N} \sum_{n=0}^{N-1} (\cos \phi + w(n)) \right]$$

From (9.18)

$$E(\hat{\phi}) = h(\underline{0}) = \arccos \left[\frac{1}{N} \sum_{n=0}^{N-1} \cos \phi \right] = \phi$$

From (9.19)

$$\text{var}(\hat{\phi}) = \left. \frac{\partial h}{\partial \underline{w}} \right|_{\underline{w}=\underline{0}}^T \sigma^2 \mathbf{I} \left. \frac{\partial h}{\partial \underline{w}} \right|_{\underline{w}=\underline{0}}$$

$$\frac{\partial h}{\partial w(i)} = - \frac{1}{\sqrt{1-u^2}} \frac{\partial u}{\partial w(i)}$$

$$\text{where } u = \frac{1}{N} \sum_{n=0}^{N-1} (\cos \phi + w(n))$$

$$\left. \frac{\partial h}{\partial w(i)} \right|_{\underline{w}=\underline{0}} = - \frac{1}{\sqrt{1-\cos^2 \phi}} \frac{1}{N} = - \frac{1}{N \sin \phi}$$

$$\text{var}(\hat{\phi}) = \sigma^2 \sum_{n=0}^{N-1} \left(\left. \frac{\partial h}{\partial w(n)} \right|_{\underline{w}=\underline{0}} \right)^2$$

$$= \sigma^2 \sum_{n=0}^{N-1} \frac{1}{N^2 \sin^2 \phi} = \frac{\sigma^2}{N \sin^2 \phi}$$

$$8) \quad \hat{\theta} = g(\underline{\tau})$$

$$\approx g(\underline{\mu}) + \sum_{k=1}^r \frac{\partial g}{\partial T_k} \Big|_{\underline{T}=\underline{\mu}} (T_k - \mu_k) \\ + \frac{1}{2} (\underline{T} - \underline{\mu})^T \frac{\partial^2 g}{\partial \underline{T} \partial \underline{T}^T} \Big|_{\underline{T}=\underline{\mu}} (\underline{T} - \underline{\mu})$$

where $\frac{\partial^2 g}{\partial \underline{T} \partial \underline{T}^T}$ is the Hessian

$$E(\hat{\theta}) = g(\underline{\mu}) + \frac{1}{2} E[(\underline{T} - \underline{\mu})^T \underline{G}(\underline{\mu}) (\underline{T} - \underline{\mu})] \\ = g(\underline{\mu}) + \frac{1}{2} E[\text{tr}((\underline{T} - \underline{\mu})^T \underline{G}(\underline{\mu}) (\underline{T} - \underline{\mu}))] \\ = g(\underline{\mu}) + \frac{1}{2} \text{tr} E[\underline{G}(\underline{\mu}) (\underline{T} - \underline{\mu}) (\underline{T} - \underline{\mu})^T] \\ = g(\underline{\mu}) + \frac{1}{2} \text{tr}(\underline{G}(\underline{\mu}) \underline{C}_T)$$

$$9) \quad \hat{\theta} \approx g(\underline{\mu}) + \frac{\partial g}{\partial \underline{T}} \Big|_{\underline{T}=\underline{\mu}}^T (\underline{T} - \underline{\mu}) \\ + \frac{1}{2} (\underline{T} - \underline{\mu})^T \underline{G}(\underline{\mu}) (\underline{T} - \underline{\mu})$$

$$\text{var}(\hat{\theta}) = E[(\hat{\theta} - E(\hat{\theta}))^2] \\ = E\left[\left(\frac{\partial g}{\partial \underline{T}} \Big|_{\underline{T}=\underline{\mu}}^T (\underline{T} - \underline{\mu}) + \frac{1}{2} (\underline{T} - \underline{\mu})^T \underline{G}(\underline{\mu}) (\underline{T} - \underline{\mu}) - \frac{1}{2} \text{tr}(\underline{G}(\underline{\mu}) \underline{C}_T)\right)^2\right]$$

using the results from Prob 9.8

But $\underline{T} - \underline{\mu} \sim N(\underline{0}, \underline{C}_T)$ so that all odd-order moments are zero. Let $\underline{x} = \underline{T} - \underline{\mu}$, $\underline{b} = \frac{\partial g}{\partial \underline{T}} \Big|_{\underline{T}=\underline{\mu}}$ and $\underline{A} = \underline{G}(\underline{\mu})$.

$$\begin{aligned} \text{var}(\hat{\theta}) &= E \left[\left(\underline{b}^T \underline{x} + \frac{1}{2} \underline{x}^T \underline{A} \underline{x} - \frac{1}{2} \text{tr}(\underline{A} \underline{C}_T) \right)^2 \right] \\ &= \underline{b}^T \underline{C}_T \underline{b} + \frac{1}{4} E[(\underline{x}^T \underline{A} \underline{x})^2] - \frac{1}{4} E(\underline{x}^T \underline{A} \underline{x}) \text{tr}(\underline{A} \underline{C}_T) \\ &\quad - \frac{1}{4} E(\underline{x}^T \underline{A} \underline{x}) \text{tr}(\underline{A} \underline{C}_T) + \frac{1}{4} \text{tr}^2(\underline{A} \underline{C}_T) \end{aligned}$$

$$\begin{aligned} \text{But } E(\underline{x}^T \underline{A} \underline{x}) &= E[\text{tr}(\underline{A} \underline{x} \underline{x}^T)] \\ &= \text{tr}(\underline{A} E(\underline{x} \underline{x}^T)) \\ &= \text{tr}(\underline{A} \underline{C}_T) \end{aligned}$$

$$\begin{aligned} E[(\underline{x}^T \underline{A} \underline{x})^2] &= \text{var}(\underline{x}^T \underline{A} \underline{x}) + E^2(\underline{x}^T \underline{A} \underline{x}) \\ &= 2 \text{tr}[(\underline{A} \underline{C}_T)^2] + \text{tr}^2(\underline{A} \underline{C}_T) \end{aligned}$$

$$\begin{aligned} \text{var}(\hat{\theta}) &= \underline{b}^T \underline{C}_T \underline{b} + \frac{1}{2} \text{tr}[(\underline{A} \underline{C}_T)^2] + \frac{1}{4} \text{tr}^2(\underline{A} \underline{C}_T) \\ &\quad - \frac{1}{2} \text{tr}^2(\underline{A} \underline{C}_T) + \frac{1}{4} \text{tr}^2(\underline{A} \underline{C}_T) \\ &= \underline{b}^T \underline{C}_T \underline{b} + \frac{1}{2} \text{tr}[(\underline{A} \underline{C}_T)^2] \end{aligned}$$

10) $g(T_1) = 1/T_1$, where $T_1 = \frac{1}{N} \sum_{n=0}^{N-1} x(n) = \bar{x}$
 $\bar{x}$ will be approximately Gaussian due to the central limit theorem

$$E(\hat{\lambda}) = g(\mu) + \frac{1}{2} \text{tr}[G(\mu) \underline{C}_T]$$

$$\text{But } \mu = E(T_1) = 1/\lambda$$

$$\underline{C}_T = \text{var}(T_1) = \text{var}(x(n))/N = \frac{1}{N\lambda^2}$$

$$G(\mu) = \frac{\partial^2 g}{\partial T^2} \bigg|_{T=\mu}$$

$$= \frac{2}{T^3} \bigg|_{T=\mu} = \frac{2}{\mu^3} = 2\lambda^3$$

$$E(\hat{\lambda}) = \lambda + \frac{1}{2} (2\lambda^3) \frac{1}{N\lambda^2} = \lambda + \lambda/N$$

$$= \lambda (1 + 1/N)$$

$$\text{var}(\hat{\lambda}) = \left(\frac{\partial g}{\partial \tau} \bigg|_{\tau=\mu} \right)^2 \text{var}(\tau)$$

$$+ \frac{1}{2} \text{tr} \left[(G(\mu) C_T)^2 \right]$$

$$= \lambda^2/N + \frac{1}{2} \left(2\lambda^3 \frac{1}{N\lambda^2} \right)^2$$

$$= \lambda^2/N + 2\lambda^2/N^2 = \frac{\lambda^2}{N} \left(1 + 2/N \right)$$

The estimator displays a bias of λ/N and an additional variance of $2\lambda^2/N^2$.

Asymptotic MLE theory is valid to first-order.

If $2/N \ll 1$, the MLE asymptotics will hold.

$$11) \quad \frac{1}{N-1} \sum_{n=0}^{N-2} A \cos(2\pi f_0 n + \phi) A \cos(2\pi f_0 (n+1) + \phi)$$

$$= \frac{A^2}{2(N-1)} \sum_{n=0}^{N-2} \left[\cos(4\pi f_0 n + 2\pi f_0 + 2\phi) \right. \\ \left. + \cos(2\pi f_0) \right]$$

As $N \rightarrow \infty$, the double-frequency term $\rightarrow 0$ for f_0 not near 0 or $1/2$. Hence, the overall expression $\rightarrow \frac{A^2}{2} \cos 2\pi f_0$.

Method of moments is valid since ensemble mean equals temporal mean as $N \rightarrow \infty$ (ergodic).

- 12) If $A \neq \sqrt{2}$, (9.20) can produce meaningless results. From Prob. 9.11, in the absence of noise, the argument of (9.20) will be approximately $A^2/2 \cos 2\pi f_0$. We need to normalize out the $A^2/2$ factor.

Now at high SNR

$$\frac{1}{N-1} \sum_{n=0}^{N-2} x(n)x(n+1) \approx \frac{1}{N-1} \sum_{n=0}^{N-2} s(n)s(n+1)$$

$$\text{where } s(n) = A \cos(2\pi f_0 n + \phi)$$

and as $N \rightarrow \infty$, this becomes $A^2/2 \cos 2\pi f_0$ from Prob. 9.11. Similarly,

$$\frac{1}{N} \sum_{n=0}^{N-1} x^2(n) \rightarrow A^2/2$$

so that $\hat{f}_0 \rightarrow f_0$. Also, from (9.18) we have $E(\hat{f}_0) = f_0$ for large N . Note that

$$\hat{f}_0 = \frac{1}{2\pi} \arccos \left[\frac{\hat{r}_{xx}(1)}{\hat{r}_{xx}(0)} \right] \text{ so that}$$

it is a method of moments estimator.

At lower SNR, as $N \rightarrow \infty$ we have

$$\hat{r}_{xx}(1) \rightarrow r_{xx}(1) = A^2/2 \cos 2\pi f_0$$

$$\hat{r}_{xx}(0) \rightarrow r_{xx}(0) = \frac{A^2}{2} + \sigma^2$$

(assuming ϕ is $U[0, 2\pi]$). Thus, as $N \rightarrow \infty$

$$\hat{f}_0 \rightarrow \frac{1}{2\pi} \arccos \left[\frac{A^2 \cos 2\pi f_0}{A^2 + 2\sigma^2} \right] \neq f_0$$

$\hat{f}_0$ will be severely biased.

$$13) \text{var}(\hat{f}_0) = \frac{\sigma^2}{(2\pi)^2 (N-1)^2 \sin^2 2\pi f_0} \cdot \left[s^2(1) + 4 \cos^2 2\pi f_0 \sum_{n=1}^{N-2} s^2(n) + s^2(N-2) \right]$$

$$f_0 = 0.25 \Rightarrow \cos 2\pi f_0 = 0$$

$$s(1) = \sqrt{2} \cos 2\pi f_0 = 0$$

$$\begin{aligned} s(N-2) &= \sqrt{2} \cos 2\pi \frac{1}{4} (N-2) \\ &= \sqrt{2} \cos \frac{\pi}{2} (N-2) = 0 \\ &\text{for } N \text{ odd} \end{aligned}$$

First-order Taylor expansion is invalid since first-order derivatives are zero. We need a second-order expansion as in Prob 9.9. To see this consider

$$\begin{aligned} f(x) &\approx f(x_0) + \left. \frac{df}{dx} \right|_{x=x_0} (x-x_0) \\ &\quad + \frac{1}{2} \left. \frac{d^2f}{dx^2} \right|_{x=x_0} (x-x_0)^2 \end{aligned}$$

Even if $(x-x_0)^2 \ll (x-x_0)$, the second-order term is not negligible if $\left. \frac{df}{dx} \right|_{x=x_0} = 0$.

Chapter 10

$$\begin{aligned}
 1) \quad \hat{\theta} &= E(\theta | \underline{x}) = \int \theta p(\theta | \underline{x}) d\theta \\
 &= \int \theta \frac{p(\underline{x} | \theta) p(\theta)}{p(\underline{x})} d\theta \\
 &= \frac{\int \theta p(\underline{x} | \theta) p(\theta) d\theta}{\int p(\underline{x} | \theta) p(\theta) d\theta} = \frac{\int \theta p(\underline{x} | \theta) \delta(\theta - \theta_0) d\theta}{\int p(\underline{x} | \theta) \delta(\theta - \theta_0) d\theta} \\
 &= \frac{\theta_0 p(\underline{x} | \theta_0)}{p(\underline{x} | \theta_0)} = \theta_0
 \end{aligned}$$

The MMSE estimator is just the true value of θ since our prior knowledge is perfect. Of course, this is not a valid estimator.

$$\begin{aligned}
 2) \quad p_{\underline{x}}(x[0], x[1] | A) &= p_{\underline{w}}(x[0] - A, x[1] - A | A) \\
 \text{But } \underline{w} &\text{ is independent of } A \\
 \Rightarrow &= p_{\underline{w}}(x[0] - A, x[1] - A) \\
 \text{And } w[0] &\text{ is independent of } w[1] \\
 \Rightarrow &= p_w(x[0] - A) p_w(x[1] - A) \\
 &= p_{x[0]}(x[0] | A) p_{x[1]}(x[1] | A)
 \end{aligned}$$

$$\underline{x} = \begin{bmatrix} A + w[0] \\ A + w[1] \end{bmatrix} = A \underline{1} + \underline{w} \sim N(0, \underline{1}\underline{1}^T + \underline{I})$$

since $A, w[0], w[1]$ are IID and $\sim N(0, 1)$

Now consider the exponent of $p(\underline{x})$ or $\underline{x}^T \underline{C}^{-1} \underline{x}$

$$\underline{C} = \begin{bmatrix} 2 & 1 \\ 1 & 2 \end{bmatrix} \Rightarrow \underline{C}^{-1} = \frac{1}{3} \begin{bmatrix} 2 & -1 \\ -1 & 2 \end{bmatrix}$$

$$\underline{x}^T \underline{C}^{-1} \underline{x} = \frac{1}{3} \underline{x}^T \begin{bmatrix} 2 & -1 \\ -1 & 2 \end{bmatrix} \underline{x} = \frac{1}{3} (2x^2[0] + 2x^2[1] - 2x[0]x[1])$$

which does not factor $\Rightarrow x[0], x[1]$ are not independent

$$3) \quad p(\underline{x}|\theta) = e^{-\sum_n (x[n] - \theta)} = e^{-N(\bar{x} - \theta)} \quad \begin{array}{l} \text{all } x[n]'s > \theta \\ \mathcal{J} = \min x[n] > \theta \end{array}$$

$$\begin{aligned} p(\theta|\underline{x}) &= \frac{p(\underline{x}|\theta) p(\theta)}{\int p(\underline{x}|\theta) p(\theta) d\theta} \\ &= \frac{e^{-N(\bar{x} - \theta)} e^{-\theta}}{\int_0^{\mathcal{J}} e^{-N(\bar{x} - \theta)} e^{-\theta} d\theta} \\ &= \frac{e^{\theta(N-1)}}{\frac{e^{\theta(N-1)}}{N-1} \Big|_0^{\mathcal{J}}} = \frac{(N-1)e^{\theta(N-1)}}{e^{(N-1)\mathcal{J}} - 1} \quad 0 < \theta < \mathcal{J} \\ &\quad \quad \quad 0 \quad \quad \quad \text{otherwise} \end{aligned}$$

$$\hat{\theta} = E[\theta|\underline{x}] = \frac{\int_0^{\mathcal{J}} \theta (N-1) e^{\theta(N-1)} d\theta}{e^{(N-1)\mathcal{J}} - 1}$$

$$= \frac{N-1}{e^{(N-1)\mathcal{J}} - 1} \left[\frac{\theta(N-1) - 1}{(N-1)^2} e^{\theta(N-1)} \Big|_0^{\mathcal{J}} \right]$$

$$\begin{aligned}
&= \frac{1}{(N-1)(e^{(N-1)\beta} - 1)} \left[(3(N-1) - 1) e^{(N-1)\beta} + 1 \right] \\
&= \frac{1}{(N-1)(e^{(N-1)\beta} - 1)} \left[1 - e^{(N-1)\beta} + (N-1)\beta e^{(N-1)\beta} \right] \\
&= \frac{3 e^{(N-1)\beta}}{e^{(N-1)\beta} - 1} - \frac{1}{N-1}
\end{aligned}$$

$$\hat{\theta} = \frac{\min x(n)}{1 - e^{-(N-1)\min x(n)}} - \frac{1}{N-1}$$

$$\begin{aligned}
4) \quad p(x|\theta) &= \frac{1}{\theta^N} \quad \text{all } x(n) \leq \theta \text{ or } \max x(n) \leq \theta \\
p(\theta) &= 1/\beta \quad 0 \leq \theta \leq \beta
\end{aligned}$$

$p(\theta|x)$ will be nonzero for $\max x(n) \leq \theta \leq \beta$
or $3 \leq \theta \leq \beta$

$$\begin{aligned}
p(\theta|x) &= \frac{p(x|\theta)p(\theta)}{\int p(x|\theta)p(\theta)d\theta} \\
&= \frac{\frac{1}{\theta^N} \frac{1}{\beta}}{\int_3^\beta \frac{1}{\theta^N} \frac{1}{\beta} d\theta} = \frac{\frac{1}{\theta^N}}{-\frac{\theta^{N-1}}{N-1} \Big|_3^\beta} \\
&= \frac{-(N-1)1/\theta^N}{\beta^{N-1} - 3^{N-1}} \\
&= \frac{1/\theta^N}{\frac{1}{(N-1)}(\beta^{-(N-1)} - 3^{-(N-1)})}
\end{aligned}$$

$$\begin{aligned}
 \hat{\theta} &= E(\theta|x) = \int \theta p(\theta|x) d\theta \\
 &= C \int_3^{\beta} \theta \frac{d\theta}{\theta^N} \quad C = \frac{1}{\frac{1}{N-1} (3^{-(N-1)} - \beta^{-(N-1)})} \\
 &= C \left. \frac{\theta^{-(N-2)}}{-(N-2)} \right|_3^{\beta} = -\frac{C}{N-2} (\beta^{-(N-2)} - 3^{-(N-2)}) \\
 &= \frac{(\max x(n))^{-(N-2)} - \beta^{-(N-2)}}{(\max x(n))^{-(N-1)} - \beta^{-(N-1)}} \cdot \frac{N-1}{N-2}
 \end{aligned}$$

Note that for β large (no prior knowledge) and N large

$$\hat{\theta} \approx \frac{N-1}{N-2} \max x(n) \approx \max x(n)$$

which agrees with the MLE.

$$\begin{aligned}
 5) \quad \text{Bmse}(\hat{\theta}) &= E_{x, \theta} \left\{ [(\theta - E(\theta|x)) + (E(\theta|x) - \hat{\theta})]^2 \right\} \\
 &= E_x \left[E_{\theta|x} \left\{ [(\theta - E(\theta|x)) + (E(\theta|x) - \hat{\theta})]^2 \right\} \right] \\
 &= E_x \left\{ E_{\theta|x} [(\theta - E(\theta|x))^2] \right. \\
 &\quad \left. + 2 E_{\theta|x} [(\theta - E(\theta|x))(E(\theta|x) - \hat{\theta})] \right. \\
 &\quad \left. + E_{\theta|x} [(E(\theta|x) - \hat{\theta})^2] \right\}
 \end{aligned}$$

The middle term is zero since conditioned on x , $\hat{\theta}$ is a constant so that

$$\begin{aligned}
& E_{\theta|x} \left[(\theta - E(\theta|x)) (E(\theta|x) - \hat{\theta}) \right] \\
&= E_{\theta|x} \left[(\theta - E(\theta|x)) (E(\theta|x) - \hat{\theta}) \right] \\
&= \underbrace{E_{\theta|x} [\theta - E(\theta|x)]}_{E(\theta|x)} (E(\theta|x) - \hat{\theta}) = 0
\end{aligned}$$

$$\text{Bmse}(\hat{\theta}) = E_x \left\{ E_{\theta|x} \left[(\theta - E(\theta|x))^2 \right] + E_{\theta|x} \left[(E(\theta|x) - \hat{\theta})^2 \right] \right\}$$

Clearly, to minimize $\text{Bmse}(\hat{\theta})$ we choose $\hat{\theta} = E(\theta|x)$ so that the last term (which is nonnegative) is zero.

b) Now A and $W(n)$ are not independent.

$$p(W(n)|A) = \frac{1}{\sqrt{2\pi\sigma_+^2}} e^{-\frac{1}{2\sigma_+^2} W^2(n)} \quad A \geq 0$$

$$\frac{1}{\sqrt{2\pi\sigma_-^2}} e^{-\frac{1}{2\sigma_-^2} W^2(n)} \quad A < 0$$

$$\Rightarrow p(X(n)|A) = \frac{1}{\sqrt{2\pi\sigma_+^2}} e^{-\frac{1}{2\sigma_+^2} (X(n)-A)^2} \quad A \geq 0$$

$$\frac{1}{\sqrt{2\pi\sigma_-^2}} e^{-\frac{1}{2\sigma_-^2} (X(n)-A)^2} \quad A < 0$$

$$\text{or } p(\underline{x}|A) = \frac{1}{(2\pi\sigma_+^2)^{N/2}} e^{-\frac{1}{2\sigma_+^2} \sum_n (X(n)-A)^2} \quad A \geq 0$$

$$\frac{1}{(2\pi\sigma_-^2)^{N/2}} e^{-\frac{1}{2\sigma_-^2} \sum_n (X(n)-A)^2} \quad A < 0$$

Since conditioned on A the $W(n)$'s and hence the $X(n)$'s are independent. Clearly,

for $\sigma_+^2 \neq \sigma_-^2$, $p(x|A) \neq p(x;A)$.
 For $\sigma_+^2 = \sigma_-^2$ they are identical.

$$7) \quad \hat{A} = \frac{\int_{-A_0}^{A_0} A \frac{1}{\sqrt{2\pi\sigma^2/N}} e^{-\frac{1}{2\sigma^2/N}(A-\bar{x})^2} dA}{\int_{-A_0}^{A_0} \frac{1}{\sqrt{2\pi\sigma^2/N}} e^{-\frac{1}{2\sigma^2/N}(A-\bar{x})^2} dA}$$

$$I_1 = \int_{-A_0}^{A_0} A e^{-a(A-\bar{x})^2} dA \quad a = \frac{1}{2\sigma^2/N}$$

$$= \int_{-A_0}^{A_0} (A-\bar{x}) e^{-a(A-\bar{x})^2} dA + \bar{x} \int_{-A_0}^{A_0} e^{-a(A-\bar{x})^2} dA$$

$$= \underbrace{\int_{-A_0-\bar{x}}^{A_0-\bar{x}} y e^{-ay^2} dy}_{I_2} + \bar{x} \int_{-A_0}^{A_0} e^{-a(A-\bar{x})^2} dA$$

$$I_2 = \left. \frac{e^{-ay^2}}{-2a} \right|_{-A_0-\bar{x}}^{A_0-\bar{x}} = -\frac{1}{2a} (e^{-a(A_0-\bar{x})^2} - e^{-a(A_0+\bar{x})^2})$$

The numerator becomes

$$= \frac{1}{\sqrt{2\pi\sigma^2/N}} \frac{1}{2} \frac{1}{\frac{1}{2\sigma^2/N}} \left(e^{-\frac{1}{2\sigma^2/N}(A_0-\bar{x})^2} - e^{-\frac{1}{2\sigma^2/N}(A_0+\bar{x})^2} \right) + \frac{\bar{x}}{\sqrt{2\pi\sigma^2/N}} \int_{-A_0}^{A_0} e^{-\frac{1}{2\sigma^2/N}(A-\bar{x})^2} dA$$

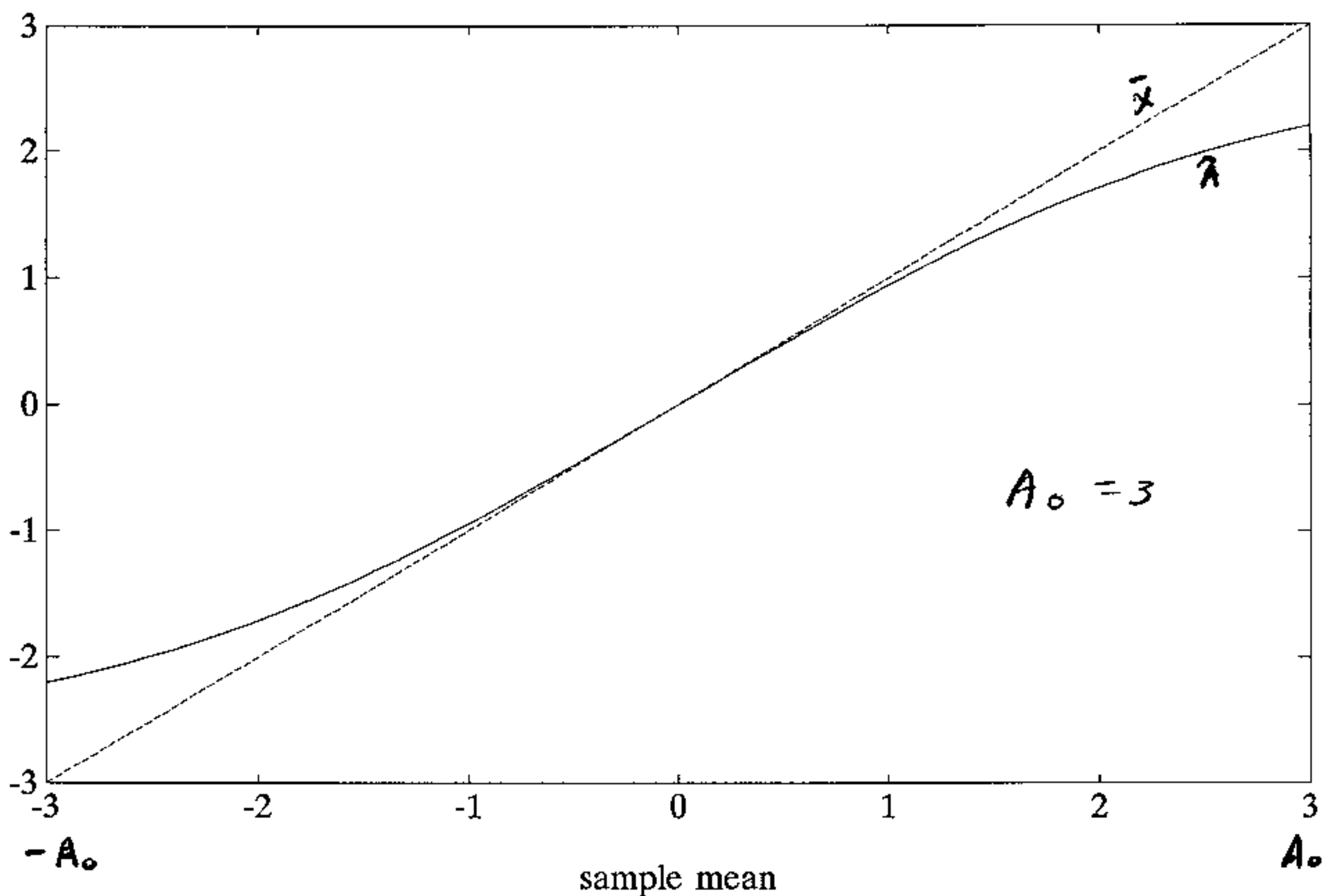
or letting $\Phi(x)$ be the CDF for a $N(0,1)$ random variable

$$\hat{A} = \bar{x} + \frac{\sqrt{\frac{\sigma^2}{N}}}{2\pi} \left[e^{-\frac{1}{2\sigma^2/N}(A_0 + \bar{x})^2} - e^{-\frac{1}{2\sigma^2/N}(A_0 - \bar{x})^2} \right]$$

$$\Phi\left(\frac{A_0 - \bar{x}}{\sqrt{\sigma^2/N}}\right) - \Phi\left(\frac{-A_0 - \bar{x}}{\sqrt{\sigma^2/N}}\right)$$

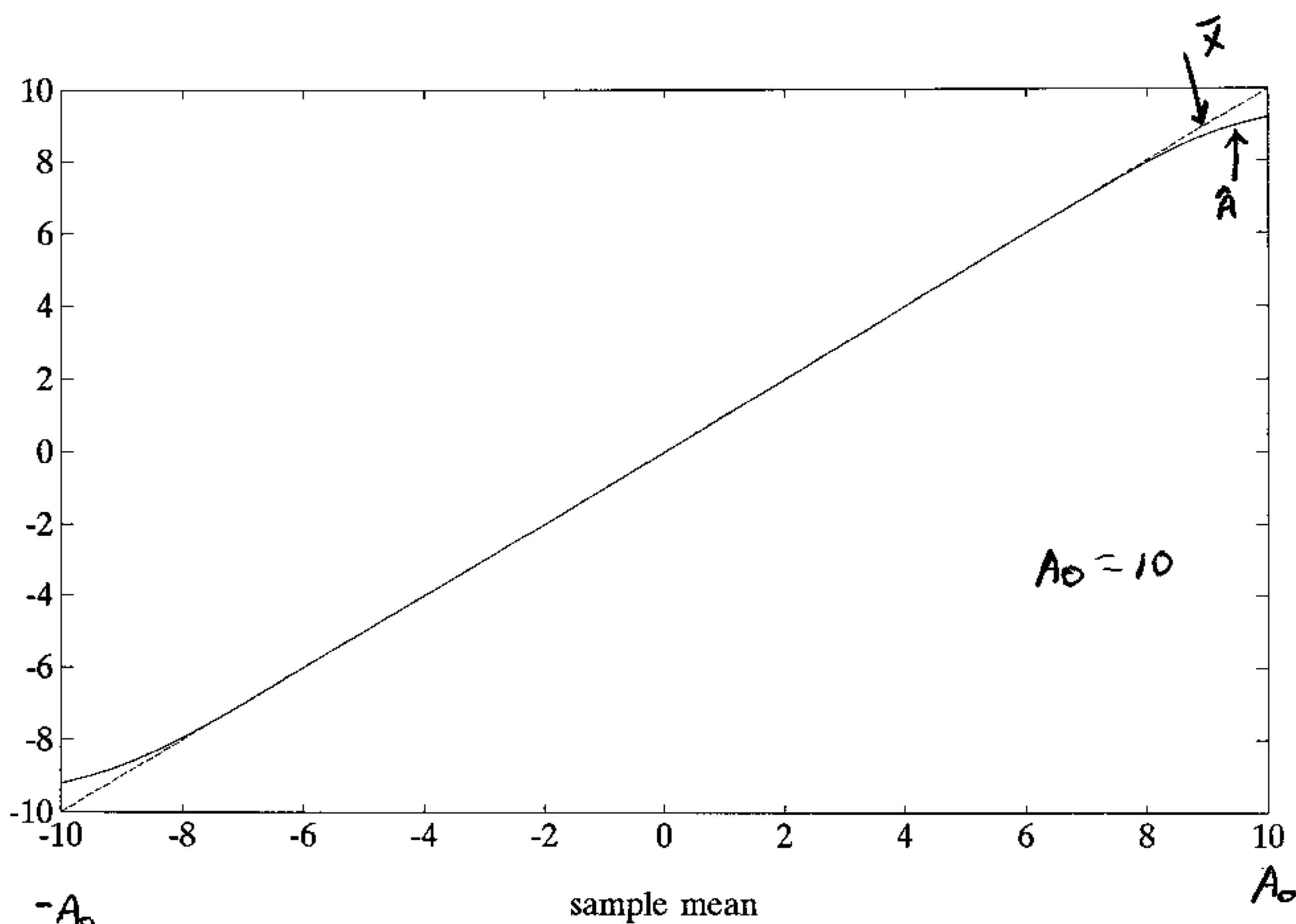
For $\sqrt{\sigma^2/N} = 1$, $A_0 = 3$ we have

$$\hat{A} = \bar{x} + \frac{e^{-\frac{1}{2}(3+\bar{x})^2} - e^{-\frac{1}{2}(3-\bar{x})^2}}{\sqrt{2\pi} \left[\Phi(3-\bar{x}) - \Phi(-3-\bar{x}) \right]}$$



and for $A_0 = 10$ (see the plot on next page). The curves are nearly identical or $\hat{A} \approx \bar{x}$ for $|\bar{x}| \leq A_0$. Also, as $\bar{x} \rightarrow \infty$,

we will always have $\hat{A} \rightarrow A_0$.



8) Let $J = E[(\theta - \hat{\theta})^2]$

$$J = E[(\theta - E(\theta) + (E(\theta) - \hat{\theta}))^2]$$

$$= E[(\theta - E(\theta))^2] + 2E[(\theta - E(\theta))(E(\theta) - \hat{\theta})] + E[(E(\theta) - \hat{\theta})^2]$$

Since $\hat{\theta}$ is a constant, the middle term is zero,

$$E[(\theta - E(\theta))(E(\theta) - \hat{\theta})] = E[(\theta - E(\theta))](E(\theta) - \hat{\theta}) \\ = (E(\theta) - E(\theta))(E(\theta) - \hat{\theta}) = 0$$

$$J = E[(\theta - E(\theta))^2] + (E(\theta) - \hat{\theta})^2 \\ \geq E[(\theta - E(\theta))^2]$$

$\Rightarrow \hat{\theta} = E(\theta)$. The minimum MSE is just $E[(\theta - E(\theta))^2] = \text{var}(\theta)$.

From Example 10.1 with no data

$$\text{Bmse}(\hat{A}) = \text{var}(\hat{A}) = \sigma_A^2$$

and with data

$$\text{Bmse}(\hat{A}) = \sigma_A^2 \frac{\sigma^2/N}{\sigma_A^2 + \sigma^2/N} < \sigma_A^2.$$

9) Prior PDF: $N(\underbrace{100}_{\mu_R}, \underbrace{0.011}_{\sigma_R^2})$

Data model: $x[n] = R + w[n]$ $n = 0, 1, \dots, N-1$

where $w[n] \sim N(0, 1)$ and $w[n]$'s

are independent

This is just Example 10.1.

$$\Rightarrow \text{Bmse}(\hat{R}) = \frac{\sigma_R^2 \sigma^2/N}{\sigma_R^2 + \sigma^2/N}$$

For the error to be 0.1 on the average

$$\Rightarrow \text{Bmse}(\hat{R}) = 0.01$$

$$0.01 = \frac{0.011/N}{0.011 + 1/N} \Rightarrow N = 9.09$$

$$\text{or } N = 10$$

Without prior knowledge or as $\sigma_R^2 \rightarrow \infty$

$$\text{Bmse}(\hat{R}) \rightarrow \sigma^2/N$$

$$0.01 = 1/N \Rightarrow \text{Require } N = 100.$$

$$10) \quad p(\theta | \underline{x}) = \frac{p(\underline{x} | \theta) p(\theta)}{\int p(\underline{x} | \theta) p(\theta) d\theta}$$

$$= \frac{\theta^N e^{-N\theta \bar{x}} \frac{\lambda^\alpha}{\Gamma(\alpha)} \theta^{\alpha-1} e^{-\lambda\theta}}{\frac{\lambda^\alpha}{\Gamma(\alpha)} \int_0^\infty \theta^{N+\alpha-1} e^{-(N\bar{x}+\lambda)\theta} d\theta}$$

$$\text{But } \int_0^\infty p(\theta) d\theta = 1 \Rightarrow \int_0^\infty \frac{\lambda^\alpha}{\Gamma(\alpha)} \theta^{\alpha-1} e^{-\lambda\theta} d\theta = 1$$

$$\int_0^\infty \theta^{\frac{N+\alpha-1}{\alpha'}} e^{-\frac{(N\bar{x}+\lambda)\theta}{\lambda'}} d\theta = \frac{\Gamma(\alpha')}{\lambda'^{\alpha'}}$$

$$= \frac{\Gamma(N+\alpha)}{(N\bar{x}+\lambda)^{N+\alpha}}$$

$$p(\theta | \underline{x}) = \frac{(N\bar{x}+\lambda)^{N+\alpha}}{\Gamma(N+\alpha)} \theta^{N+\alpha-1} e^{-(N\bar{x}+\lambda)\theta}$$

0

 $\theta > 0$ $\theta < 0$

This PDF is also a Gamma PDF with parameters $\alpha' = N + \alpha$, $\lambda' = N\bar{x} + \lambda$.

Only the PDF parameters change from the prior to the posterior PDF. Note that the trick here is to have $p(\underline{x} | \theta)$ and $p(\theta)$ of the same form and to retain it after multiplication. The denominator is just a

scaling constant.

- 11) The scatter diagram in Figure 10.9a indicates that the PDF of the random vector $\begin{bmatrix} h \\ w \end{bmatrix}$ is concentrated within an ellipse. This could be a bivariate Gaussian PDF. Hence, assuming $\begin{bmatrix} h \\ w \end{bmatrix}$ is Gaussian, we estimate w based on h using a MMSE estimator or from (10.16)

$$\hat{w} = E(w) + \frac{\text{cov}(h, w)}{\text{var}(h)} (h - E(h))$$

From Fig 10.9b there does not appear to be any correlation between height and weight
 $\Rightarrow \text{Cov}(h, w) = 0$ or $\hat{w} = E(w) = 150$
 for any height.

$$12) p(x, y) = \frac{1}{2\pi \det^{1/2}(\underline{C})} e^{-\frac{1}{2} \underline{x}^T \underline{C}^{-1} \underline{x}}$$

$$g(y) = p(x_0, y) = \frac{1}{2\pi \det^{1/2}(\underline{C})} e^{-\frac{1}{2} h(y)}$$

$$\text{where } h(y) = \begin{bmatrix} x_0 \\ y \end{bmatrix}^T \begin{bmatrix} 1 & \rho \\ \rho & 1 \end{bmatrix}^{-1} \begin{bmatrix} x_0 \\ y \end{bmatrix}$$

$$= \frac{\begin{bmatrix} x_0 \\ y \end{bmatrix}^T \begin{bmatrix} 1 & -\rho \\ -\rho & 1 \end{bmatrix} \begin{bmatrix} x_0 \\ y \end{bmatrix}}{1 - \rho^2}$$

$$= (x_0^2 + y^2 - 2\rho x_0 y) / (1 - \rho^2)$$

$g(y)$ is maximized when $h(y)$ is minimized or

$$\frac{dh}{dy} = \frac{2y - 2\rho x_0}{1 - \rho^2} = 0 \Rightarrow y = \rho x_0$$

Now from (10.16) with $E(X) = E(y) = 0$,

$$\text{Cov}(X, y) = \rho, \text{Var}(X) = 1$$

$$\Rightarrow E(y|X) = \rho X$$

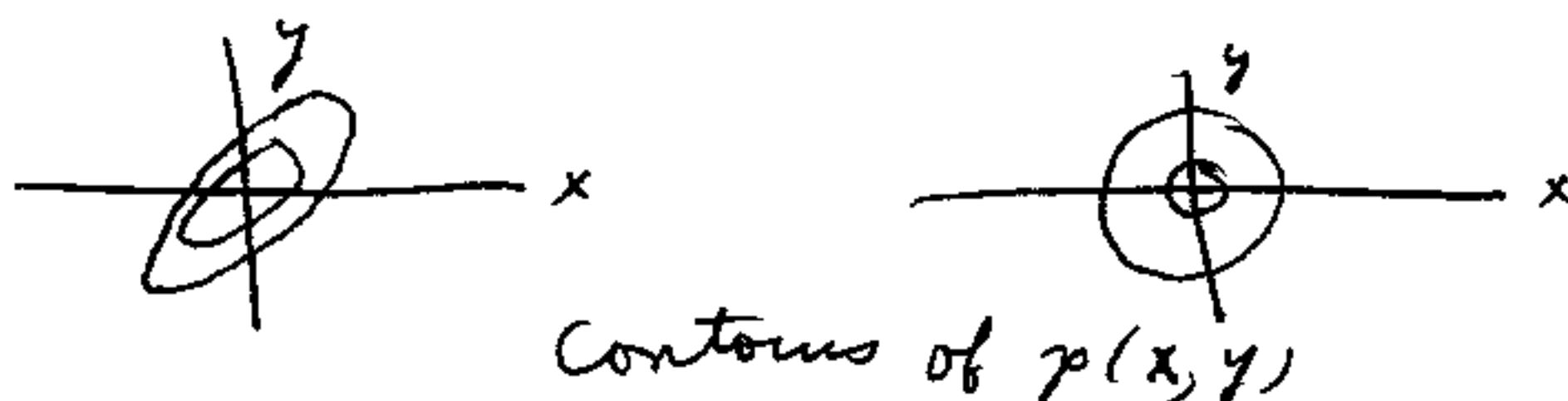
The joint PDF $p(x_0, y)$ has the identical form as $p(y|x_0)$, with the only difference being the normalization since

$$p(y|x_0) = \frac{p(x_0, y)}{\int p(x_0, y) dy}$$

Hence, for a given x_0 the maximum of $p(y|x_0)$ is the same as that for $p(x_0, y)$.

Also, the posterior PDF is Gaussian so that the mode (maximizing value of y) is identical to the mean.

$$\text{If } \rho = 0, E(y|x_0) = 0 = E(y)$$



13) Since

$$(\underline{A} + \underline{B}\underline{C}\underline{D})^{-1} = \underline{A}^{-1} - \underline{A}^{-1}\underline{B}(\underline{D}\underline{A}^{-1}\underline{B} + \underline{C}^{-1})^{-1}\underline{D}\underline{A}^{-1}$$

$$(\underline{C}_0^{-1} + \underline{H}^T \underline{C}_W^{-1} \underline{H})^{-1} = \underline{C}_0 - \underline{C}_0 \underline{H}^T (\underline{H} \underline{C}_0 \underline{H}^T + \underline{C}_W)^{-1} \underline{H} \underline{C}_0$$

$\Rightarrow (10.33)$

Now, using the inversion lemma again

$$\begin{aligned} (\underline{C}_0^{-1} + \underline{H}^T \underline{C}_W^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}_W^{-1} &= \underline{C}_0 \underline{H}^T \underline{C}_W^{-1} \\ &\quad - \underline{C}_0 \underline{H}^T (\underline{H} \underline{C}_0 \underline{H}^T + \underline{C}_W)^{-1} \underline{H} \underline{C}_0 \underline{H}^T \underline{C}_W^{-1} \end{aligned}$$

$$= \underline{C}_0 \underline{H}^T \left[\underline{C}_W^{-1} - (\underline{H} \underline{C}_0 \underline{H}^T + \underline{C}_W)^{-1} \underline{H} \underline{C}_0 \underline{H}^T \underline{C}_W^{-1} \right]$$

$$\underline{E} = \underline{A}^{-1} - (\underline{B} + \underline{A})^{-1} \underline{B} \underline{A}^{-1}$$

$$\begin{aligned} \underline{E} &= \underline{A}^{-1} - \underline{A}^{-1} \underline{B} (\underline{B}^{-1} \underline{A}) (\underline{B} + \underline{A})^{-1} \underline{B} \underline{A}^{-1} \\ &= \underline{A}^{-1} - \underline{A}^{-1} \underline{B} \left[\underline{B}^{-1} (\underline{B} + \underline{A}) (\underline{B}^{-1} \underline{A})^{-1} \right]^{-1} \underline{A}^{-1} \\ &= \underline{A}^{-1} - \underline{A}^{-1} \underline{B} \left[\underline{B}^{-1} (\underline{B} + \underline{A}) (\underline{A}^{-1} \underline{B}) \right]^{-1} \underline{A}^{-1} \\ &= \underline{A}^{-1} - \underline{A}^{-1} \underline{B} \left[\underline{I} + \underline{A}^{-1} \underline{B} \right]^{-1} \underline{A}^{-1} \\ &= (\underline{A} + \underline{B})^{-1} \quad \text{by inversion lemma} \end{aligned}$$

$$\Rightarrow \underline{E} = (\underline{H} \underline{C}_0 \underline{H}^T + \underline{C}_W)^{-1}$$

14) This is the Bayesian linear model with

$$\underline{H} = [1, r, \dots, r^{N-1}]^T$$

Use (10.32), (10.33) since $\underline{H}$ is a column vector and thus $\underline{H}^T \underline{C}_W^{-1} \underline{H}$ is a scalar.

No matrix inversion required as in (10.28)

$$\begin{aligned}
 \hat{A} &= 0 + \left(\frac{1}{\sigma_A^2} + \frac{H^T H}{\sigma^2} \right)^{-1} H^T \frac{1}{\sigma^2} (\underline{x} - \underline{0}) \\
 &= \frac{\frac{1}{\sigma^2} \sum_{n=0}^{N-1} x[n] r^n}{\frac{1}{\sigma_A^2} + \frac{\sum_{n=0}^{N-1} r^{2n}}{\sigma^2}} \\
 &= \frac{\sigma_A^2}{\sigma_A^2 \sum_{n=0}^{N-1} r^{2n} + \sigma^2} \sum_{n=0}^{N-1} x[n] r^n
 \end{aligned}$$

From (10.13)

$$\text{Bmse}(\hat{A}) = \int \text{var}(A|\underline{x}) p(\underline{x}) d\underline{x}$$

$$\text{But } \text{var}(\hat{A}) = \frac{1}{\frac{1}{\sigma_A^2} + \frac{1}{\sigma^2} \sum_{n=0}^{N-1} r^{2n}} \text{ from (10.33)}$$

which does not depend on $\underline{x}$

$$\Rightarrow \text{Bmse}(\hat{A}) = \frac{1}{\frac{1}{\sigma_A^2} + \frac{1}{\sigma^2} \sum_{n=0}^{N-1} r^{2n}}$$

$$15) \ln p(\theta) = \ln \frac{1}{\sqrt{2\pi\sigma_\theta^2}} - \frac{1}{2\sigma_\theta^2} (\theta - \mu_\theta)^2$$

$$E[\ln p(\theta)] = -\frac{1}{2} \ln 2\pi\sigma_\theta^2 - \frac{1}{2\sigma_\theta^2} \underbrace{E[(\theta - \mu_\theta)^2]}_{\sigma_\theta^2}$$

$$= -\frac{1}{2} [1 + \ln 2\pi\sigma_\theta^2]$$

$$H(\theta) = \frac{1}{2} (1 + \ln 2\pi\sigma_\theta^2)$$

The more concentrated the PDF or for σ_θ^2 small,

the smaller will be $H(\theta)$. Higher entropy
 $\Rightarrow$ larger σ_θ^2 or a more random θ .

$$\begin{aligned}
 I &= H(\theta) - H(\theta|x) \\
 &= -\int p(\theta) \ln p(\theta) d\theta + \iint p(x, \theta) \ln p(\theta|x) dx d\theta \\
 &= -\iint p(x, \theta) dx \ln p(\theta) d\theta + \quad " \\
 &= -\iint p(x, \theta) \ln p(\theta) dx d\theta \\
 &\quad + \iint p(x, \theta) \ln p(\theta|x) dx d\theta \\
 &= \iint p(x, \theta) \ln \frac{p(\theta|x)}{p(\theta)} dx d\theta \\
 &= \underbrace{\iint p(\theta|x) \ln \frac{p(\theta|x)}{p(\theta)} d\theta p(x) dx}_{\geq 0 \text{ by given inequality}} \\
 &\geq 0 \quad \text{since } p(x) \geq 0. \\
 &= 0 \quad \text{if and only if } p(\theta|x) = p(\theta)
 \end{aligned}$$

Makes sense since if posterior PDF is the same as prior PDF $\Rightarrow$ no information.

16) $H(\theta) = \frac{1}{2} (1 + \ln 2\pi \sigma_\theta^2)$ from Prob 10.15.
 Similarly, since $p(\theta|x)$ is Gaussian,
 $H(\theta|x) = \frac{1}{2} (1 + \ln 2\pi \sigma_{\theta|x}^2)$
 $\Rightarrow I = H(\theta) - H(\theta|x) = \frac{1}{2} \ln \sigma_\theta^2 / \sigma_{\theta|x}^2$

$$\text{But } \sigma_\theta^2 = \sigma_A^2, \quad \sigma_{\theta|x}^2 = \sigma_{A|x}^2 \\ = \frac{\sigma_A^2 \sigma^2/N}{\sigma_A^2 + \sigma^2/N}$$

$$\begin{aligned} \mathcal{I} &= \frac{1}{2} \ln \frac{\sigma_A^2 + \sigma^2/N}{\sigma^2/N} \\ &= \frac{1}{2} \ln \left(1 + \frac{\sigma_A^2}{\sigma^2/N} \right) \end{aligned}$$

- 17) From Prob 10.16 we want to maximize $\mathcal{I}$. We do so by letting $\sigma_A^2 \rightarrow \infty$. It doesn't matter what we choose for μ_A . This choice swamps out the prior as we have already observed.

Chapter 11

- 1) This is just the DC level in WGN or Example 10.1.
From (10.11), the MMSE estimator is

$$\hat{u} = \frac{\sigma_0^2}{\sigma_0^2 + \sigma^2/N} \bar{x} + \frac{\sigma^2/N}{\sigma_0^2 + \sigma^2/N} \mu_0$$

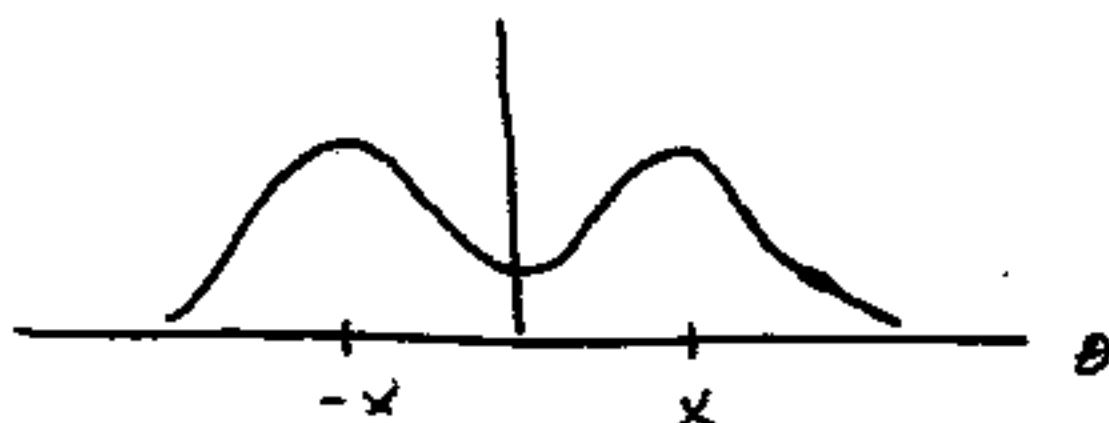
Also, since the posterior PDF is Gaussian

$\Rightarrow$ MMSE estimator = MAP estimator

As $\sigma_0^2 \rightarrow 0$, $\hat{u} \rightarrow \mu_0$ prior knowledge dominates

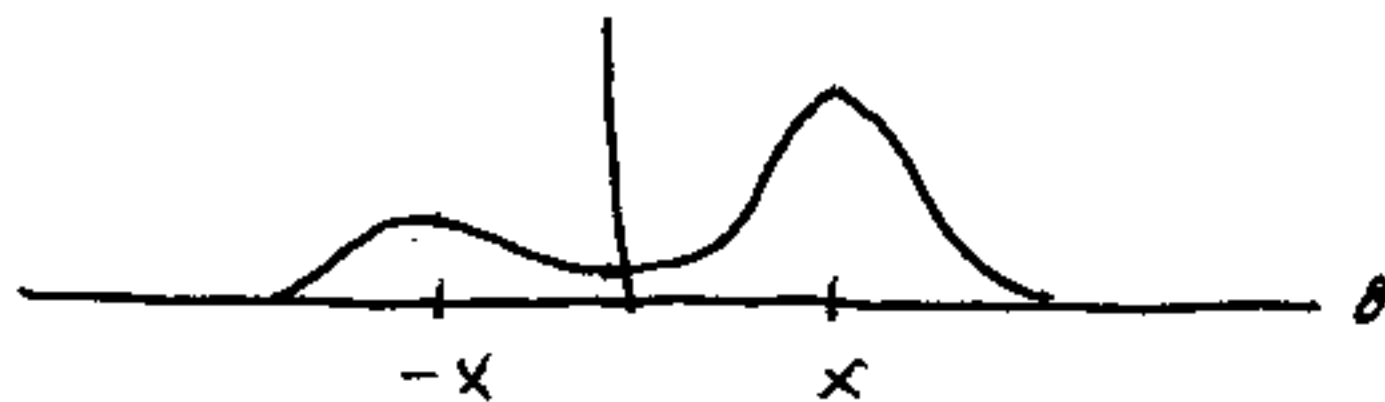
As $\sigma_0^2 \rightarrow \infty$, $\hat{u} \rightarrow \bar{x}$ data " "

2)



$$\epsilon = \frac{1}{2}$$

MMSE estimator is $E(\theta|x) = 0$ due to even symmetry of PDF. MAP estimator is mode or $\pm x$ (not unique)



$$\epsilon = \frac{3}{4}$$

To find MMSE estimator

$$\begin{aligned} \hat{\theta} = E(\theta|x) &= \int_{-\infty}^{\infty} \theta \frac{\epsilon}{\sqrt{2\pi}} e^{-\frac{1}{2}(\theta-x)^2} d\theta \\ &+ \int_{-\infty}^{\infty} \theta \frac{1-\epsilon}{\sqrt{2\pi}} e^{-\frac{1}{2}(\theta+x)^2} d\theta \end{aligned}$$

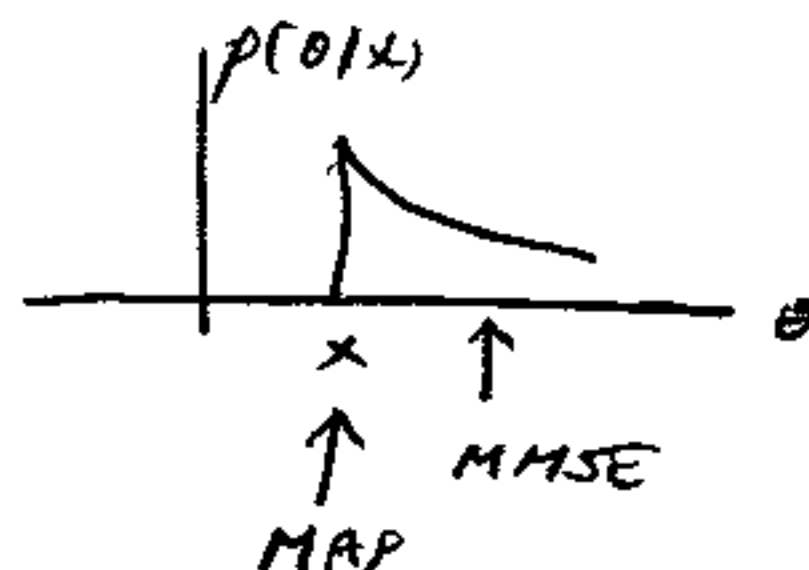
$$= \epsilon x + (1-\epsilon)(-x) = x(2\epsilon-1) = x/2$$

The MAP estimator is $\hat{\theta} = x$. Note that $\text{MAP} \neq \text{MMSE}$.

3) MMSE:

$$\begin{aligned}\hat{\theta} &= \int_x^\infty \theta e^{-(\theta-x)} d\theta = e^x \left[-\theta e^{-\theta} - e^{-\theta} \right] \Big|_x^\infty \\ &= e^x (x e^{-x} + e^{-x}) = x+1\end{aligned}$$

MAP est. is just $\hat{\theta} = x$



$$\begin{aligned}4) \quad g(A) &= p(x|A) p(A) \\ &= \frac{1}{(2\pi\sigma^2)^{N/2}} e^{-\frac{1}{2\sigma^2} \sum_n (x(n)-A)^2} \lambda e^{-\lambda A} \quad A > 0\end{aligned}$$

0

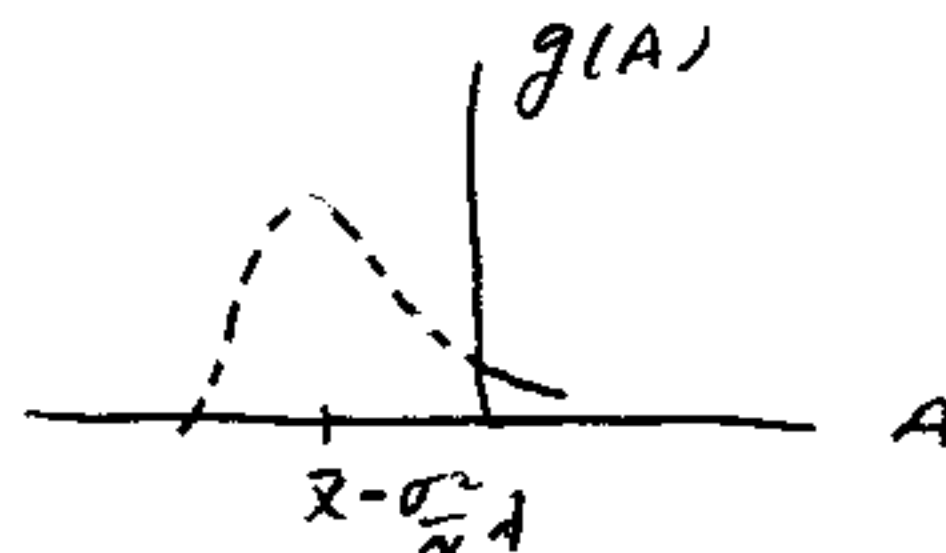
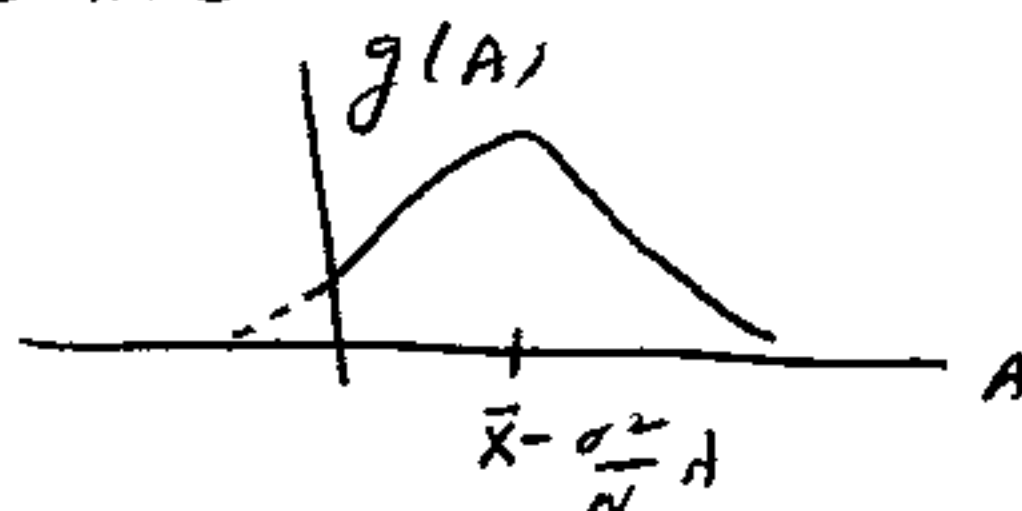
$A < 0$

$$\frac{\partial \ln g}{\partial A} = \frac{1}{\sigma^2} \sum_n (x(n)-A) - \lambda = 0$$

$$\Rightarrow \hat{A} = \bar{x} - \frac{\sigma^2}{N} \lambda$$

Note that if $\hat{A}$ is negative we use $\hat{A} = 0$

Since



$$\Rightarrow \hat{A} = \max(0, \bar{x} - \frac{\sigma^2}{N} \lambda)$$

5) From (11.17)

$$\hat{\underline{\theta}} = E[\underline{\theta} | \underline{x}]$$

$$= E[\underline{\theta}] + \underline{C}_{\theta x} \underline{C}_{xx}^{-1} (\underline{x} - E(\underline{x}))$$

From (11.27)

$$Bmse(\hat{\theta}_i) = [M\hat{\theta}]_{ii}$$

$$= [\underline{C}_{\theta\theta} - \underline{C}_{\theta x} \underline{C}_{xx}^{-1} \underline{C}_{x\theta}]_{ii}$$

$$\text{If } \underline{C}_{\theta x} = 0, \quad \hat{\underline{\theta}} = E[\underline{\theta}]$$

$$Bmse(\hat{\theta}_i) = [\underline{C}_{\theta\theta}]_{ii} = \text{var}(\theta_i)$$

Data is irrelevant to estimator since $\underline{\theta}$ and $\underline{x}$ are independent (due to Gaussian assumption)

6) The Jacobian is

$$\underline{J} = \frac{\partial [A \phi]^T}{\partial (a \ b)^T} = \begin{bmatrix} \partial A / \partial a & \partial A / \partial b \\ \partial \phi / \partial a & \partial \phi / \partial b \end{bmatrix}$$

$$= \begin{bmatrix} \frac{a}{\sqrt{a^2 + b^2}} & \frac{b}{\sqrt{a^2 + b^2}} \\ \frac{b/a^2}{1 + (b/a)^2} & \frac{-1/a}{1 + (b/a)^2} \end{bmatrix}$$

$$= \begin{bmatrix} a/A & b/A \\ b/A^2 & -a/A^2 \end{bmatrix}$$

$$|\det \underline{J}| = \left| -a^2/A^3 - b^2/A^3 \right| = 1/A$$

$$p(A, \phi) = \frac{p(a, b)}{|\det \underline{J}|}$$

Note: transformation is 1-1

$$= \frac{1}{2\pi\sigma_0^2} e^{-\frac{1}{2\sigma_0^2}(a^2+b^2)} \frac{1}{A}$$

$$= \frac{A}{\sigma_0^2} e^{-\frac{1}{2\sigma_0^2}A^2} \cdot \frac{1}{2\pi} \quad \begin{matrix} A > 0 \\ 0 \leq \phi \leq 2\pi \end{matrix}$$

$$p(A) = \int_0^{2\pi} p(A, \phi) d\phi = \frac{A}{\sigma_0^2} e^{-\frac{1}{2\sigma_0^2}A^2}$$

$$p(\phi) = \int_0^\infty p(A, \phi) dA = \frac{1}{2\pi}$$

Since $p(A, \phi)$ factors, A and ϕ are independent.

- 7) For Bayesian linear model, MMSE estimator = MAP estimator since $p(\underline{\theta}|\underline{x})$ is Gaussian. But MAP estimator maximizes $p(\underline{x}|\underline{\theta})p(\underline{\theta})$. With no prior information this is equivalent to maximizing $p(\underline{x}|\underline{\theta})$. In the Bayesian model $p(\underline{x}|\underline{\theta}) = p(\underline{x}; \underline{\theta})$. Thus, maximizing $p(\underline{x}; \underline{\theta})$, which yields the MLE or MVU estimator, also yields the MMSE estimator.

8)
$$\begin{aligned} \hat{\underline{\theta}}[n] &= E(\underline{\theta}[n]|\underline{x}) = E(\underline{A}\underline{\theta}[n-1]|\underline{x}) \\ &= \underline{A} E(\underline{\theta}[n-1]|\underline{x}) = \underline{A} \hat{\underline{\theta}}[n-1] \end{aligned}$$
 due to linearity of the expectation
 Let $n=1 \Rightarrow \hat{\underline{\theta}}[1] = \underline{A} \hat{\underline{\theta}}[0]$
 $n=2 \Rightarrow \hat{\underline{\theta}}[2] = \underline{A} \hat{\underline{\theta}}[1] = \underline{A}^2 \hat{\underline{\theta}}[0]$
 etc.

$$9) \quad \underline{\theta}(n) = \begin{bmatrix} x(n) \\ y(n) \\ v_x \\ v_y \end{bmatrix} = \begin{bmatrix} x(n-1) + v_x \\ y(n-1) + v_y \\ v_x \\ v_y \end{bmatrix}$$

$$\begin{aligned} \text{Since } x(n) - x(n-1) &= x(0) + v_x n \\ &= x(0) - v_x(n-1) \\ &= v_x \end{aligned}$$

And similarly for the y component.

$$\underline{\theta}(n) = \underbrace{\begin{bmatrix} 1 & 0 & 1 & 0 \\ 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}}_{\underline{A}} \underbrace{\begin{bmatrix} x(n-1) \\ y(n-1) \\ v_x \\ v_y \end{bmatrix}}_{\underline{\theta}(n-1)}$$

$$\Rightarrow \text{from Prob. 11.8 } \hat{\underline{\theta}}(n) = \underline{A}^n \hat{\underline{\theta}}(0)$$

10) $\hat{\theta} = \frac{1}{\bar{x} + N_N}$ As $N \rightarrow \infty$, $\hat{\theta} \rightarrow 1/\bar{x}$ and thus as $N \rightarrow \infty$ for a given realization of θ , $\bar{x} \rightarrow E(x) = 1/\theta \Rightarrow \hat{\theta} \rightarrow \theta$. In general, as $N \rightarrow \infty$ the MAP estimator is just the value that maximizes $p(\underline{x}|\theta)$ or the Bayesian MLE. For a given realization of θ , if we assume $p(\underline{x}|\theta) = p(\underline{x}; \theta)$, then we can treat $\hat{\theta}$ as an MLE. (The family of PDFs is now characterized by $p(\underline{x}|\theta)$). If the MLE is consistent, then so will be the MAP estimator.

$$\begin{aligned}
 11) \quad R &= E[C(\underline{\epsilon})] \\
 &= \iint C(\underline{\epsilon}) p(\underline{x}, \underline{\theta}) d\underline{x} d\underline{\theta} \\
 &= \int \underbrace{\int C(\underline{\epsilon}) p(\underline{\theta}|\underline{x}) d\underline{\theta}}_{\text{minimize over } \hat{\underline{\theta}}} p(\underline{x}) d\underline{x}
 \end{aligned}$$

$$\begin{aligned}
 \int C(\underline{\epsilon}) p(\underline{\theta}|\underline{x}) d\underline{\theta} &= \int_{\{\underline{\theta}: \|\underline{\epsilon}\| > \delta\}} p(\underline{\theta}|\underline{x}) d\underline{\theta} \\
 &= \int_{\{\underline{\theta}: \|\underline{\theta} - \hat{\underline{\theta}}\| > \delta\}} p(\underline{\theta}|\underline{x}) d\underline{\theta} = \int_{\{\underline{\theta}: \|\underline{\theta} - \hat{\underline{\theta}}\| \leq \delta\}^c} p(\underline{\theta}|\underline{x}) d\underline{\theta} \quad c \text{ denotes Complement}
 \end{aligned}$$

$$= 1 - \int_{\{\underline{\theta}: \|\underline{\theta} - \hat{\underline{\theta}}\| \leq \delta\}} p(\underline{\theta}|\underline{x}) d\underline{\theta}$$

As $\delta \rightarrow 0$, this is minimized if the integral is maximized or if we choose $\hat{\underline{\theta}} = \arg \max_{\underline{\theta}} p(\underline{\theta}|\underline{x})$

12) MAP estimator of $\underline{\theta}$ maximizes

$$\begin{aligned}
 p(\underline{x}|\underline{\theta}) p(\underline{\theta}) &= p(\underline{x}, \underline{\theta}) \\
 \text{If } \underline{\alpha} &= \underline{A} \underline{\theta}, \quad \frac{\partial \underline{\alpha}}{\partial \underline{\theta}} = \underline{A}
 \end{aligned}$$

$$p(\underline{x}, \underline{\alpha}) = \frac{p(\underline{x}, \underline{\theta})}{\left| \det \frac{\partial \underline{\alpha}}{\partial \underline{\theta}} \right|} = \frac{p(\underline{x}, \underline{\theta})}{|\det \underline{A}|}$$

But $\underline{A}$ does not depend on $\underline{\alpha}$ and $\underline{\theta} = \underline{A}^{-1} \underline{\alpha}$

$$\text{so that } p(\underline{x}, \underline{\alpha}) = \frac{p_{\underline{x}, \underline{\theta}}(\underline{x}, \underline{A}^{-1} \underline{\alpha})}{|\det \underline{A}|}$$

The MAP estimator of $\underline{\alpha}$ maximizes $p_{\underline{x}, \underline{\theta}}(\underline{x}, \underline{A}^{-1}\underline{\alpha})$ or because $\underline{0} = \underline{A}^{-1}\underline{\alpha}$ is invertible we can maximize $p(\underline{x}, \underline{0}) \Rightarrow$ maximizing value is $\hat{\underline{0}}$ and since $\underline{\alpha} = \underline{A}\underline{0} \Rightarrow \hat{\underline{\alpha}} = \underline{A}\hat{\underline{0}}$.

$$13) \quad \underline{x}^T \underline{C}^{-1} \underline{x} = \underline{x}^T \underline{D}^T \underline{D} \underline{x} = \underline{y}^T \underline{y} \quad \text{where } \underline{y} = \underline{D} \underline{x}$$

$$\text{But } \underline{C}_{\underline{y}} = E(\underline{y} \underline{y}^T) = E(\underline{D} \underline{x} \underline{x}^T \underline{D}^T) = \underline{D} \underline{C} \underline{D}^T \\ = \underline{D} (\underline{D}^T \underline{D})^{-1} \underline{D}^T = \underline{I}$$

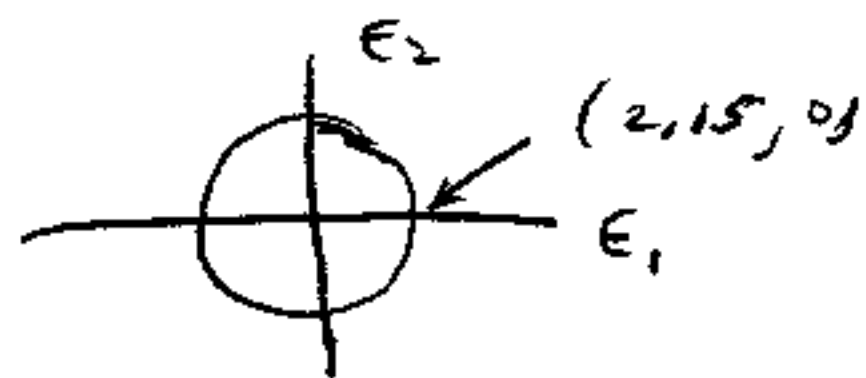
$$\Rightarrow \underline{y} \sim N(\underline{0}, \underline{I}) \quad \text{and}$$

$$\underline{y}^T \underline{y} = y_1^2 + y_2^2 \sim \chi_2^2$$

$$\text{since } \left. \begin{array}{l} y_1 \sim N(0, 1) \\ y_2 \sim N(0, 1) \end{array} \right\} \text{ independent}$$

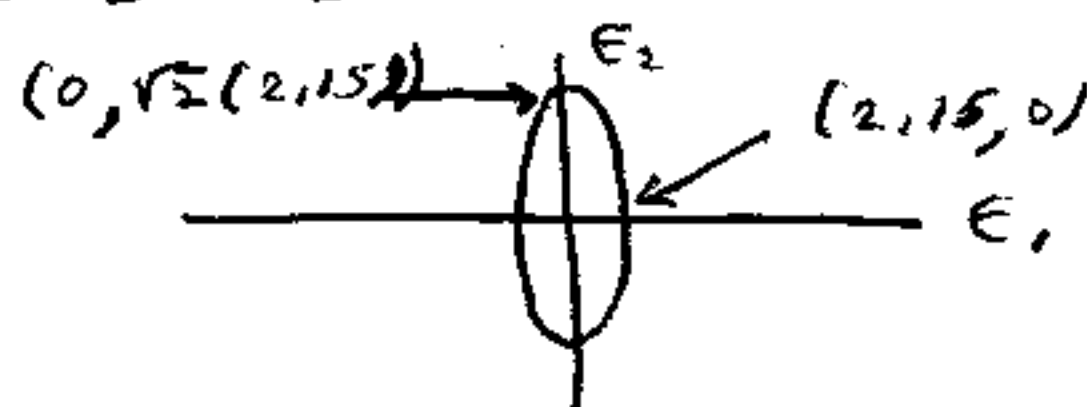
$$14) \quad \underline{\epsilon}^T \underline{M} \hat{\underline{\sigma}}^{-1} \underline{\epsilon} = 2 \ln \frac{1}{1-p} \\ = (2.15)^2$$

$$a) \quad \underline{\epsilon}^T \underline{\epsilon} = (2.15)^2$$



$$b) \quad \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix}^{-1} = \begin{bmatrix} 1 & 0 \\ 0 & 1/2 \end{bmatrix}$$

$$\underline{\epsilon}^T \underline{M} \hat{\underline{\sigma}}^{-1} \underline{\epsilon} = \epsilon_1^2 + \frac{1}{2} \epsilon_2^2 = (2.15)^2$$



$$c) \quad \underline{M}_{\hat{\theta}}^{-1} = \frac{1}{3} \begin{pmatrix} 2 & -1 \\ -1 & 2 \end{pmatrix}$$

$$[\epsilon_1, \epsilon_2] \frac{1}{3} \begin{pmatrix} 2 & -1 \\ -1 & 2 \end{pmatrix} \begin{pmatrix} \epsilon_1 \\ \epsilon_2 \end{pmatrix} = (2,15)^2$$

$$2\epsilon_1^2 + 2\epsilon_2^2 - 2\epsilon_1\epsilon_2 = 3(2,15)^2$$

$$\epsilon_1^2 - \epsilon_1\epsilon_2 + \epsilon_2^2 = \frac{3}{2}(2,15)^2$$

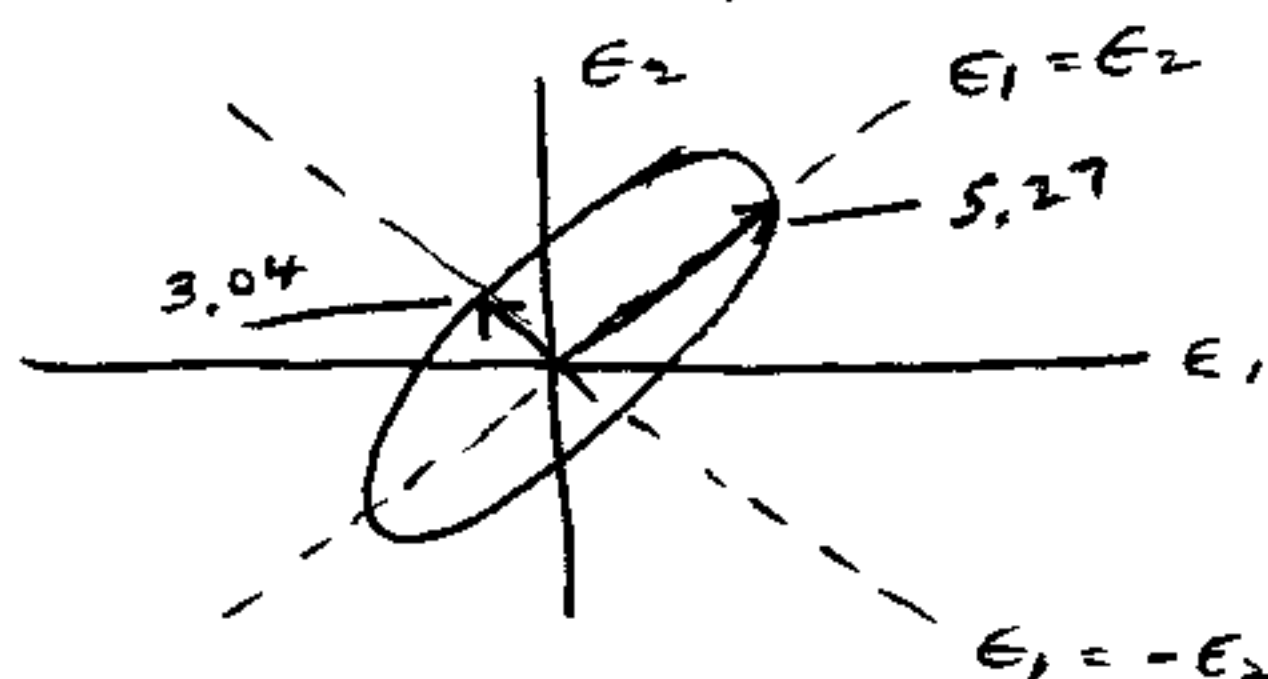
This is a rotated ellipse. It can be rewritten in standard form as

$$\frac{(\epsilon_1 + \epsilon_2)^2}{4} + \frac{(\epsilon_1 - \epsilon_2)^2}{4/3} = \frac{3}{2}(2,15)^2$$

$$\text{or} \quad \frac{(\epsilon_1 + \epsilon_2)^2}{a^2} + \frac{(\epsilon_1 - \epsilon_2)^2}{b^2} = 1$$

$$\text{where } a = 5,27$$

$$b = 3,04$$



$$15) \quad \underline{M}_{\hat{\theta}} = (\underline{C}\hat{\theta}' + H^T \underline{C}\hat{w}' H)^{-1} = (\underline{C}\hat{\theta}' + \underline{C}\hat{w}')^{-1}$$

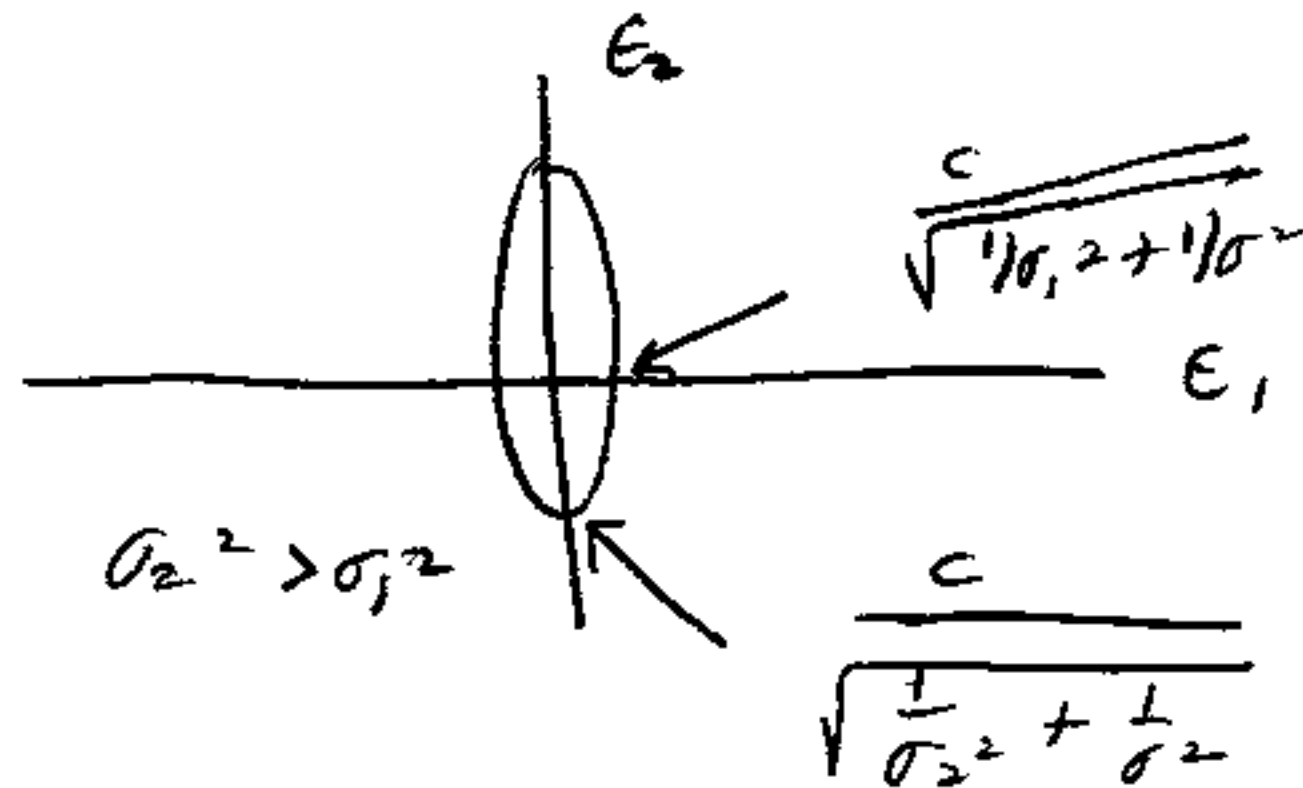
$$\underline{M}_{\hat{\theta}}^{-1} = \underline{C}\hat{\theta}' + \underline{C}\hat{w}' = \begin{pmatrix} 1/\sigma_1^2 + 1/\sigma^2 & 0 \\ 0 & 1/\sigma_2^2 + 1/\sigma^2 \end{pmatrix}$$

$$\epsilon^T \underline{M}_{\hat{\theta}}^{-1} \epsilon = 2 \ln \frac{1}{1-p}$$

$$\epsilon_1^2 \left(\frac{1}{\sigma_1^2} + \frac{1}{\sigma^2} \right) + \epsilon_2^2 \left(\frac{1}{\sigma_2^2} + \frac{1}{\sigma^2} \right) = 2 \ln \frac{1}{1-p}$$

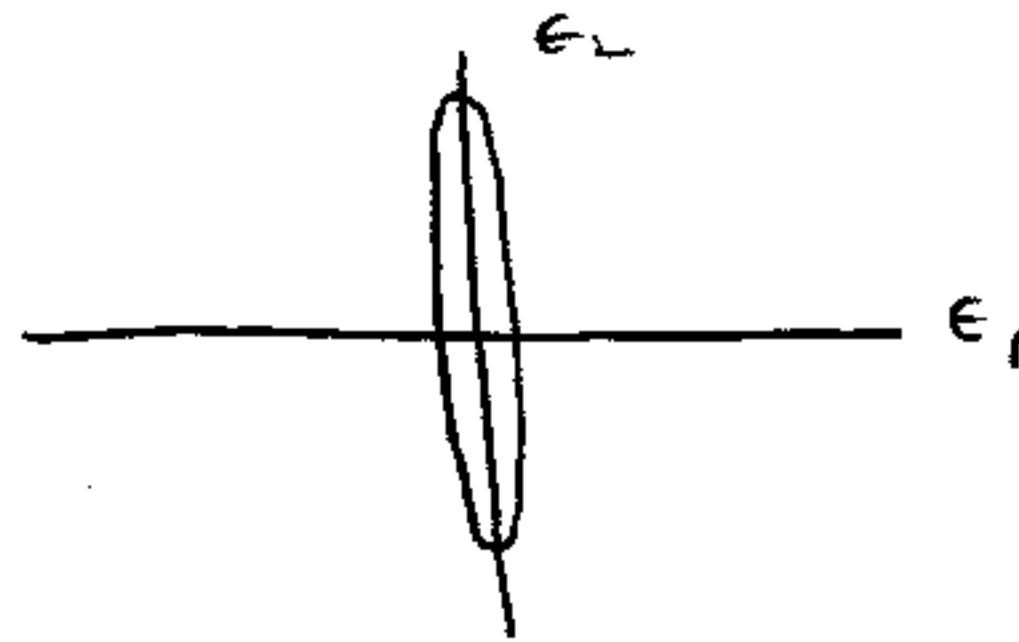
If $\sigma_1^2 = \sigma_2^2$, we have a circle.

If $\sigma_2^2 > \sigma_1^2$, we have an ellipse



$$C = \sqrt{2 \ln \frac{1}{1-p}}$$

For $\sigma_2^2 \gg \sigma_1^2$ the ellipse appears as



Most of uncertainty is along E_2 direction, as expected.

16) This is the Bayesian linear model.
From (11.33)

$$\hat{\theta} = \underline{\mu}_\theta + (\underline{C}_\theta^{-1} + \underline{H}^T \underline{C}_w^{-1} \underline{H})^{-1} \underline{H}^T \underline{C}_w^{-1} (\underline{x} - \underline{H} \underline{\mu}_\theta)$$

where $\underline{\mu}_\theta = [A_0 \ B_0]^T$

$$\underline{C}_\theta = \begin{pmatrix} \sigma_A^2 & 0 \\ 0 & \sigma_B^2 \end{pmatrix}$$

$$\underline{H} = \begin{bmatrix} 1 & -M \\ \vdots & -M+1 \\ \vdots & \vdots \\ 1 & M \end{bmatrix}$$

$$\underline{C}_w = \sigma^2 \underline{I}$$

$$\Rightarrow \underline{H}^T \underline{C}_w^{-1} \underline{H} = \frac{1}{\sigma^2} \underline{H}^T \underline{H} = \frac{1}{\sigma^2} \begin{bmatrix} N & 0 \\ 0 & \sum n^2 \end{bmatrix} \quad \text{where } N = 2M+1$$

$$(\underline{C}_\theta^{-1} + \underline{H}^T \underline{C}_w^{-1} \underline{H})^{-1} = \begin{bmatrix} \frac{1}{\sigma_A^2} + \frac{N}{\sigma^2} & 0 \\ 0 & \frac{1}{\sigma_B^2} + \frac{\sum n^2}{\sigma^2} \end{bmatrix}^{-1}$$

$$= \begin{bmatrix} \frac{1}{\frac{1}{\sigma_A^2} + \frac{N}{\sigma^2}} & 0 \\ 0 & \frac{1}{\frac{1}{\sigma_B^2} + \frac{\sum n^2}{\sigma^2}} \end{bmatrix}$$

$$\underline{H}^T \underline{C} \underline{W}^{-1} (\underline{X} - \underline{H} \underline{M}_0) = \frac{1}{\sigma^2} (\underline{H}^T \underline{X} - \underline{H}^T \underline{H} \underline{M}_0)$$

$$\underline{H}^T \underline{H} = \begin{bmatrix} N & 0 \\ 0 & \sum n^2 \end{bmatrix}$$

$$\underline{H}^T \underline{C} \underline{W}^{-1} (\underline{X} - \underline{H} \underline{M}_0) = \frac{1}{\sigma^2} \begin{bmatrix} \sum x(n) - N A_0 \\ \sum n x(n) - \sum n^2 B_0 \end{bmatrix}$$

$$\hat{A} = A_0 + \frac{\frac{1}{\sigma^2} \left[\sum x(n) - N A_0 \right]}{\frac{1}{\sigma_A^2} + \frac{N}{\sigma^2}}$$

$$= A_0 + \frac{N / \sigma^2}{\frac{1}{\sigma_A^2} + \frac{N}{\sigma^2}} (\bar{x} - A_0)$$

$$\hat{B} = B_0 + \frac{\frac{1}{\sigma^2} \left[\sum n x(n) - \sum n^2 B_0 \right]}{\frac{1}{\sigma_B^2} + \frac{\sum n^2}{\sigma^2}}$$

$$= B_0 + \frac{\frac{\sum_{n=-M}^M n^2 / \sigma^2}{\frac{1}{\sigma_B^2} + \frac{\sum_{n=-M}^M n^2 / \sigma^2}} \left[\frac{\sum_{n=-M}^M n x(n)}{\sum_{n=-M}^M n^2} - B_0 \right]$$

$$\begin{aligned} \text{Bmse}(\hat{A}) &= \left[(\underline{C} \underline{\theta}^{-1} + \underline{H}^T \underline{C} \underline{W}^{-1} \underline{H})^{-1} \right]_{11} \\ &= \left(\frac{1}{\sigma_A^2} + \frac{N}{\sigma^2} \right)^{-1} \end{aligned}$$

$$\text{Bmse}(\hat{\underline{B}}) = ((\underline{C}_0^{-1} + \underline{H}^T \underline{C}_W^{-1} \underline{H})^{-1})_{22}$$

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

The intercept ($\hat{A}$) will benefit most from prior knowledge since the reduction in Bmse due to the data is much larger for the slope ($\sum n^2 / \sigma^2$) than for the intercept (N / σ^2).

$$17) \quad \underline{S} = \begin{bmatrix} 1 & -M \\ 1 & -M+1 \\ \vdots & \vdots \\ 1 & M \end{bmatrix} \underline{\theta} = \underline{H} \underline{\theta}$$

$$\hat{\underline{S}} = E(\underline{S} | \underline{x}) = E(\underline{H} \underline{\theta} | \underline{x}) = \underline{H} \hat{\underline{\theta}}$$

where $\hat{\underline{\theta}}$ is given in Prob 10.16.

$$\underline{E} = \underline{S} - \hat{\underline{S}} = \underline{S} - \underline{H} \hat{\underline{\theta}} = \underline{H} (\underline{\theta} - \hat{\underline{\theta}})$$

$$E(\underline{E}) = \underline{H} \{E(\underline{\theta}) - E(\hat{\underline{\theta}})\}$$

$$\text{But } E(\hat{\underline{\theta}}) = \underline{\mu}_\theta \Rightarrow E(\underline{E}) = \underline{H} (\underline{\mu}_\theta - \underline{\mu}_\theta) = \underline{0}$$

The covariance is

$$E(\underline{E} \underline{E}^T) = \underline{H} E[(\underline{\theta} - \hat{\underline{\theta}})(\underline{\theta} - \hat{\underline{\theta}})^T] \underline{H}^T$$

$$= \underline{H} \underline{M}_\theta \underline{H}^T$$

$$= \begin{bmatrix} 1 & -M \\ \vdots & \vdots \\ 1 & M \end{bmatrix} \begin{bmatrix} \frac{1}{\sigma_A^2 + N/\sigma^2} & 0 \\ 0 & \frac{1}{\sigma_B^2 + \frac{\sum n^2}{\sigma^2}} \end{bmatrix} \begin{bmatrix} 1 \dots 1 \\ -M \dots M \end{bmatrix}$$

$$= \frac{1}{\sigma_A^2 + \frac{N}{\sigma^2}} \underline{1} \underline{1}^T + \frac{1}{\sigma_B^2 + \frac{\sum n^2}{\sigma^2}} \underline{m} \underline{m}^T$$

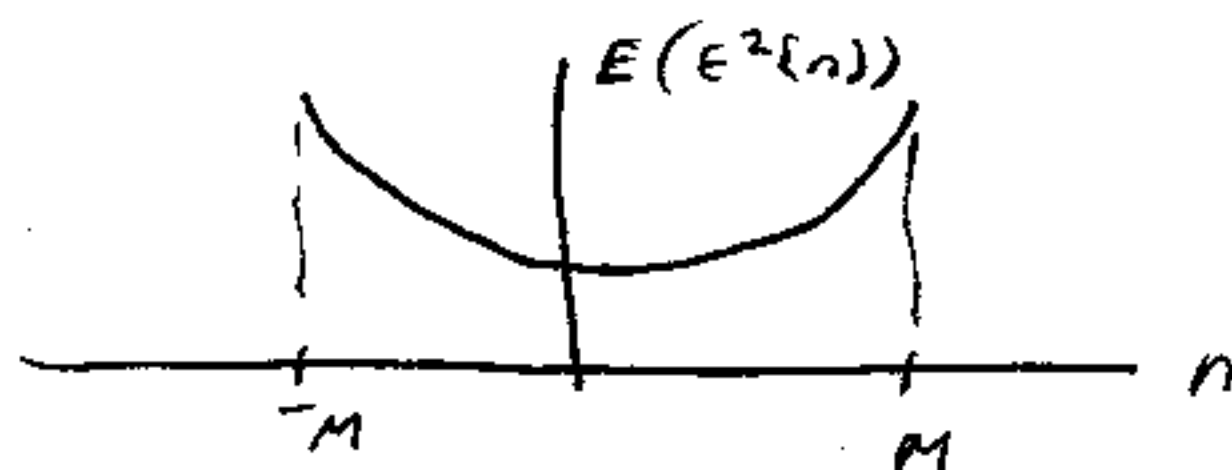
where $\underline{m} = [-M \dots M]^T$

Since $\underline{x}$, $\underline{\theta}$ are jointly Gaussian, $\underline{\epsilon}$ is also Gaussian. Note, however, that $\underline{C}_{\epsilon}$ is singular, being of rank 2.

The mean squared error is for $\hat{s}(n)$:

$$(\underline{C}_{\epsilon})_{n+M+1, n+M+1}$$

$$= \frac{1}{\frac{1}{\sigma_A^2} + \frac{N}{\sigma^2}} + \frac{1}{\frac{1}{\sigma_B^2} + \frac{\sum n^2}{\sigma^2}} n^2$$



As we depart from $n=0$ the signal estimation error increases. This is because any errors in $\underline{B}$ are magnified by n due to the signal dependence being Bn .

18) From (11.40)

$$\begin{aligned} \hat{\underline{s}} &= \underline{C}_s (\underline{C}_s + \sigma^2 \underline{I})^{-1} \underline{x} \\ &= [(\underline{C}_s + \sigma^2 \underline{I}) \underline{C}_s^{-1}]^{-1} \underline{x} \\ &= (\underline{I} + \sigma^2 \underline{C}_s^{-1})^{-1} \underline{x} \end{aligned}$$

$$(\underline{I} + \sigma^2 \underline{C}_s^{-1}) \hat{\underline{s}} = \underline{x}$$

$$(\underline{C}_s + \sigma^2 \underline{I}) \hat{\underline{s}} = \underline{C}_s \underline{x}$$

Let $\underline{x} = [x[-M] \dots x[M]]^T$

$\underline{s} = [s[-M] \dots s[M]]^T$ so that

$$[\underline{r}_s]_{ij} = E[s(i)s(j)] \quad i, j = -M, \dots, 0, \dots, M$$

$$= r_{ss}(i-j)$$

$$\Rightarrow \sum_{j=-M}^M r_{xx}(i-j) \hat{s}(j) = \sum_{j=-M}^M r_{ss}(i-j) x(j)$$

for $i = -M, \dots, M$

As $M \rightarrow \infty \Rightarrow$

$$\sum_{j=-\infty}^{\infty} r_{xx}(i-j) \hat{s}(j) = \sum_{j=-\infty}^{\infty} r_{ss}(i-j) x(j)$$

$$-\infty < i < \infty$$

or $r_{xx}(n) \star \hat{s}(n) = r_{ss}(n) \star x(n)$

Taking Fourier transforms

$$P_{xx}(f) \hat{S}(f) = P_{ss}(f) X(f)$$

$$\Rightarrow H(f) = \frac{P_{ss}(f)}{P_{xx}(f)} = \frac{P_{ss}(f)}{P_{ss}(f) + \sigma^2}$$

Wiener filter emphasizes high SNR regions
and attenuates low SNR regions since

$$H(f) = \frac{\eta(f)}{\eta(f) + 1} \quad \text{where } \eta(f) = \frac{P_{ss}(f)}{\sigma^2}$$

If $\eta(f) \gg 1$ (high SNR), $H(f) \approx 1$

If $\eta(f) \ll 1$ (low SNR), $H(f) \approx 0$.

$$E(x^3) = 0$$

$$E(x^4) = \int_{-\frac{1}{2}}^{\frac{1}{2}} x^4 dx = \frac{1}{80}$$

$$\begin{bmatrix} \frac{1}{80} & 0 & \frac{1}{12} \\ 0 & \frac{1}{12} & 0 \\ \frac{1}{12} & 0 & 1 \end{bmatrix} \begin{pmatrix} a \\ b \\ c \end{pmatrix} = \begin{pmatrix} -\frac{1}{2\pi^2} \\ 0 \\ 0 \end{pmatrix}$$

$$\Rightarrow \hat{a} = -90/\pi^2$$

$$\hat{b} = 0$$

$$\hat{c} = 15/2\pi^2$$

$$\text{MSE}(\hat{\theta}) = \frac{1}{2} - (-90/\pi^2)(-15/2\pi^2) - 0 - 0 = 0.04$$

Using a linear estimator we have

$$\hat{\theta} = b x(0) + c$$

$$\begin{bmatrix} E(x^2) & E(x) \\ E(x) & 1 \end{bmatrix} \begin{pmatrix} b \\ c \end{pmatrix} = \begin{pmatrix} E(\theta x) \\ E(\theta) \end{pmatrix}$$

$$\text{But } E(\theta x) = 0, E(\theta) = 0 \Rightarrow \hat{b} = \hat{c} = 0$$

and $\hat{\theta} = 0$. The minimum MSE is
 $E(\theta^2) = 1/2$.

2) From (12.27)

$$\hat{A} = \mu_A + \left(1/\sigma_A^2 + \frac{\underline{h}^T \underline{h}}{\sigma^2} \right)^{-1} \frac{\underline{h}^T}{\sigma^2} (\underline{x} - \underline{h} \mu_A)$$

$$\text{where } \underline{h} = [1, r, \dots, r^{N-1}]^T$$

$$= \mu_A + \frac{\sum_{n=0}^{N-1} r^n (x(n) - r^n \mu_A)}{\underbrace{\sigma_A^2}_{\sigma_A^2} + \sum_{n=0}^{N-1} r^{2n}}$$

$$\text{Bmse}(\hat{A}) = \frac{1}{\frac{1}{\sigma_A^2} + \frac{1}{\sigma^2} \sum_{n=0}^{N-1} r^{2n}}$$

from (12.29), (12.30).

3) $\hat{x}_2 = a, x_1 + b$

Let $\theta = x_2$. Then from (12.6)

$$\begin{aligned} \hat{\theta} &= E(\theta) + C_{\theta x_1} C_{x_1 x_1}^{-1} (x_1 - E(x_1)) \\ &= E(x_2) + \frac{E(x_1 x_2)}{E(x_1^2)} (x_1 - E(x_1)) \end{aligned}$$

$$\begin{aligned} \text{Bmse}(\hat{\theta}) &= C_{\theta\theta} - C_{\theta x_1} C_{x_1 x_1}^{-1} C_{x_1 \theta} \quad \text{from (12.8)} \\ &= \frac{E(x_2^2) - E^2(x_1 x_2)}{E(x_1^2)} \end{aligned}$$

$$\text{But } C_{xx} = \begin{bmatrix} E(x_1^2) & E(x_1 x_2) \\ E(x_2 x_1) & E(x_2^2) \end{bmatrix}$$

This is singular if and only if

$$E(x_1^2)E(x_2^2) - E^2(x_1 x_2) = 0$$

and is equivalent to $\text{Bmse}(\hat{\theta}) = 0$.

In general, if

$$\hat{x}_1 = \sum_{i=2}^N a_i x_i$$

$$\text{Bmse}(\hat{x}_1) = E((x_1 - \hat{x}_1)^2)$$

$$\begin{aligned}
&= E \left[\left(x_1 - \sum_{i=2}^N a_i x_i \right)^2 \right] \\
&= E \left[\left(\underline{b}^T \underline{x} \right)^2 \right] \quad \text{where } b_1 = 1 \\
&\quad \quad \quad b_i = -a_i \quad i=2, \dots, N \\
&= \underline{b}^T \underline{C}_{xx} \underline{b}
\end{aligned}$$

But $\underline{C}_{xx}$ is positive semidefinite and for the BMSE to be zero we require $\underline{b}^T \underline{C}_{xx} \underline{b} = 0$ for some $\underline{b} \neq \underline{0}$. Hence, $\underline{C}_{xx}$ must be singular.

4) $(x, x) = E(x^2) \geq 0$ and $= 0$ if and only if $x = 0$

$(x, y) = (y, x)$ obvious

$$\begin{aligned}
(c_1 x_1 + c_2 x_2, y) &= E((c_1 x_1 + c_2 x_2) y) \\
&= E(c_1 x_1 y + c_2 x_2 y) \\
&= c_1 E(x_1 y) + c_2 E(x_2 y) \\
&= c_1 (x_1, y) + c_2 (x_2, y)
\end{aligned}$$

5) $(x, x) = 0 \Rightarrow \text{cov}(x, x) = 0 \Rightarrow \text{var}(x) = 0$
 $\nRightarrow x = 0$ but that x is a constant

6) Using (12.20)

$$\hat{\underline{s}} = \underline{C}_{sx} \underline{C}_{xx}^{-1} \underline{x}$$

$$\underline{C}_{xx} = E(\underline{x} \underline{x}^T) = \underline{R}_{ss} + \underline{R}_{ww} = (\sigma_s^2 + \sigma^2) \underline{I}$$

$$\begin{aligned}
\underline{C}_{sx} &= E(\underline{s} \underline{x}^T) = E(\underline{s} (\underline{s} + \underline{w})^T) \\
&= E(\underline{s} \underline{s}^T) = \sigma_s^2 \underline{I}
\end{aligned}$$

$$\hat{\underline{s}} = \frac{\sigma_s^2}{\sigma_s^2 + \sigma^2} \underline{x} \quad \text{or} \quad \hat{s}(n) = \frac{\sigma_s^2}{\sigma_s^2 + \sigma^2} x(n)$$

From (12.21)

$$\begin{aligned}
 \underline{M}_{\hat{S}} &= \underline{C}_{SS} - \underline{C}_{SX} (\underline{X}\underline{X}'^{-1}) \underline{C}_{XS} \\
 &= \sigma_S^2 \underline{I} - \frac{(\sigma_S^2)^2}{\sigma_S^2 + \sigma^2} \underline{I} \\
 &= \sigma_S^2 \left[1 - \frac{\sigma_S^2}{\sigma_S^2 + \sigma^2} \right] \underline{I} \\
 &= \frac{\sigma_S^2 \sigma^2}{\sigma_S^2 + \sigma^2} \underline{I}
 \end{aligned}$$

$$\begin{aligned}
 7) \quad \hat{\underline{\theta}} &= E(\underline{\theta}) + \underline{C}_{\theta X} \underline{C}_{XX}^{-1} (\underline{x} - E(\underline{x})) \\
 \underline{M}_{\hat{\theta}} &= E[(\underline{\theta} - \hat{\underline{\theta}})(\underline{\theta} - \hat{\underline{\theta}})^T] \\
 &= E\left[(\underline{\theta} - E(\underline{\theta}) - \underline{C}_{\theta X} \underline{C}_{XX}^{-1} (\underline{x} - E(\underline{x}))) \right. \\
 &\quad \left. (\underline{\theta} - E(\underline{\theta}) - \underline{C}_{\theta X} \underline{C}_{XX}^{-1} (\underline{x} - E(\underline{x})))^T\right] \\
 &= \underline{C}_{\theta\theta} - E[(\underline{\theta} - E(\underline{\theta}))(\underline{x} - E(\underline{x}))^T] \underline{C}_{XX}^{-1} \underline{C}_{X\theta} \\
 &\quad - \underline{C}_{\theta X} \underline{C}_{XX}^{-1} E[(\underline{x} - E(\underline{x}))(\underline{\theta} - E(\underline{\theta}))^T] \\
 &\quad + \underline{C}_{\theta X} \underline{C}_{XX}^{-1} \underline{C}_{XX} \underline{C}_{XX}^{-1} \underline{C}_{X\theta} \\
 &= \underline{C}_{\theta\theta} - \underline{C}_{\theta X} \underline{C}_{XX}^{-1} \underline{C}_{X\theta} - \underline{C}_{\theta X} \underline{C}_{XX}^{-1} \underline{C}_{X\theta} \\
 &\quad + \underline{C}_{\theta X} \underline{C}_{XX}^{-1} \underline{C}_{X\theta} \\
 &= \underline{C}_{\theta\theta} - \underline{C}_{\theta X} \underline{C}_{XX}^{-1} \underline{C}_{X\theta}
 \end{aligned}$$

$$BMS_e(\hat{\theta}_i) = E[(\theta_i - \hat{\theta}_i)^2]$$

where the expectation is with respect to $p(\underline{x}, \theta_i)$

$$\begin{aligned}
 \text{But } BMS_e(\hat{\theta}_i) &= \int (\theta_i - \hat{\theta}_i)^2 p(\underline{x}, \theta_i) d\theta_i d\underline{x} \\
 &= \int \dots \int (\theta_i - \hat{\theta}_i)^2 p(\underline{x}, \underline{\theta}) d\underline{\theta} d\underline{x}
 \end{aligned}$$

since $\hat{\theta}_i$ depends on $\underline{x}$ only

$$\begin{aligned}
 &= \int \cdots \int [(\underline{\theta} - \hat{\underline{\theta}})(\underline{\theta} - \hat{\underline{\theta}})^T]_{ii} p(\underline{x}, \underline{\theta}) d\underline{\theta} d\underline{x} \\
 &= \left[E \left\{ (\underline{\theta} - \hat{\underline{\theta}})(\underline{\theta} - \hat{\underline{\theta}})^T \right\} \right]_{ii} = [\underline{M} \hat{\underline{\theta}}]_{ii}
 \end{aligned}$$

8) $\hat{\underline{\alpha}} = E(\underline{\alpha}) + \underline{C}_{\alpha x} \underline{C}_{xx}^{-1} (\underline{x} - E(\underline{x}))$
 But $E(\underline{\alpha}) = \underline{A} E(\underline{\theta}) + \underline{b}$
 $\underline{C}_{\alpha x} = E \{ (\underline{\alpha} - E(\underline{\alpha})) (\underline{x} - E(\underline{x}))^T \}$
 $= E \{ \underline{A} (\underline{\theta} - E(\underline{\theta})) (\underline{x} - E(\underline{x}))^T \}$
 $= \underline{A} \underline{C}_{\theta x}$
 $\Rightarrow \hat{\underline{\alpha}} = \underline{A} E(\underline{\theta}) + \underline{b} + \underline{A} \underline{C}_{\theta x} \underline{C}_{xx}^{-1} (\underline{x} - E(\underline{x}))$
 $= \underline{A} \hat{\underline{\theta}} + \underline{b}$

$$\begin{aligned}
 \hat{\underline{\alpha}} &= E(\underline{\alpha}) + \underline{C}_{\alpha x} \underline{C}_{xx}^{-1} (\underline{x} - E(\underline{x})) \\
 \text{If } \underline{\alpha} &= \underline{\theta}_1 + \underline{\theta}_2, \\
 E(\underline{\alpha}) &= E(\underline{\theta}_1) + E(\underline{\theta}_2) \\
 \underline{C}_{\alpha x} &= E \{ (\underline{\theta}_1 + \underline{\theta}_2 - E(\underline{\theta}_1) - E(\underline{\theta}_2)) (\underline{x} - E(\underline{x}))^T \} \\
 &= E \{ (\underline{\theta}_1 - E(\underline{\theta}_1)) (\underline{x} - E(\underline{x}))^T \} \\
 &\quad + E \{ (\underline{\theta}_2 - E(\underline{\theta}_2)) (\underline{x} - E(\underline{x}))^T \} \\
 &= \underline{C}_{\theta_1 x} + \underline{C}_{\theta_2 x} \\
 \Rightarrow \hat{\underline{\alpha}} &= E(\underline{\theta}_1) + E(\underline{\theta}_2) + (\underline{C}_{\theta_1 x} + \underline{C}_{\theta_2 x}) \underline{C}_{xx}^{-1} (\underline{x} - E(\underline{x})) \\
 &= \hat{\underline{\theta}}_1 + \hat{\underline{\theta}}_2
 \end{aligned}$$

9) $\hat{A}(N-1) = \frac{\sigma_A^2}{\sigma_A^2 + \sigma^2/N} \bar{x} + \frac{\sigma^2/N}{\sigma_A^2 + \sigma^2/N} \mu_A$

$$\hat{A}(N) = \frac{N\sigma_A^2}{(N+1)\sigma_A^2 + \sigma^2} \frac{1}{N} \left(\sum_{n=0}^{N-1} x[n] + x[N] \right)$$

$$\begin{aligned}
& + \frac{\sigma^2}{(N+1)\sigma_A^2 + \sigma^2} \mu_A \\
= & \frac{N\sigma_A^2}{(N+1)\sigma_A^2 + \sigma^2} \frac{\sigma_A^2 + \sigma^2/N}{\sigma_A^2} \underbrace{\frac{\sigma_A^2}{\sigma_A^2 + \sigma^2/N} \frac{1}{N} \sum_{n=0}^{N-1} x(n)}_{\hat{A}(N-1)} \\
& + \frac{\sigma_A^2}{(N+1)\sigma_A^2 + \sigma^2} x(N) + \underbrace{\frac{N\sigma_A^2 + \sigma^2}{(N+1)\sigma_A^2 + \sigma^2} \frac{\sigma^2}{N\sigma_A^2 + \sigma^2} \mu_A}_{\hat{A}(N-1)} \\
= & \frac{N\sigma_A^2 + \sigma^2}{(N+1)\sigma_A^2 + \sigma^2} \hat{A}(N) + \frac{\sigma_A^2}{(N+1)\sigma_A^2 + \sigma^2} x(N) \\
= & \hat{A}(N-1) + \left(\frac{N\sigma_A^2 + \sigma^2}{(N+1)\sigma_A^2 + \sigma^2} - 1 \right) \hat{A}(N-1) \\
& + \frac{\sigma_A^2}{(N+1)\sigma_A^2 + \sigma^2} x(N) \\
= & \hat{A}(N-1) - \frac{\sigma_A^2}{(N+1)\sigma_A^2 + \sigma^2} \hat{A}(N-1) \\
& + \frac{\sigma_A^2}{(N+1)\sigma_A^2 + \sigma^2} x(N) \\
= & \hat{A}(N-1) + \frac{\sigma_A^2}{(N+1)\sigma_A^2 + \sigma^2} (x(N) - \hat{A}(N-1))
\end{aligned}$$

Since $\text{BMS}(\hat{A}(N-1))$ for the case $\mu_A \neq 0$ is also given by (12.31), the rest of the derivation is identical. Thus, we have the identical set of equations.

10) This procedure is illustrated in Figure 12.5.

$$\underline{e}_1 = \frac{\begin{bmatrix} 1 \\ 1 \\ 2 \end{bmatrix}}{\sqrt{1+1+4}} = \begin{bmatrix} 1/\sqrt{6} \\ 1/\sqrt{6} \\ 2/\sqrt{6} \end{bmatrix}$$

$$\begin{aligned} \underline{z}_2 &= \underline{x}_2 - (\underline{x}_2, \underline{e}_1) \underline{e}_1 = \begin{bmatrix} 1 \\ 0 \\ 1 \end{bmatrix} - [1 \ 0 \ 1] \begin{bmatrix} 1/\sqrt{6} \\ 1/\sqrt{6} \\ 2/\sqrt{6} \end{bmatrix} \begin{bmatrix} 1/\sqrt{6} \\ 1/\sqrt{6} \\ 2/\sqrt{6} \end{bmatrix} \\ &= \begin{bmatrix} 1 \\ 0 \\ 1 \end{bmatrix} - \frac{3}{\sqrt{6}} \begin{bmatrix} 1/\sqrt{6} \\ 1/\sqrt{6} \\ 2/\sqrt{6} \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \\ 1 \end{bmatrix} - \begin{bmatrix} 1/2 \\ 1/2 \\ 1 \end{bmatrix} = \begin{bmatrix} 1/2 \\ -1/2 \\ 0 \end{bmatrix} \end{aligned}$$

$$\underline{e}_2 = \frac{\begin{bmatrix} 1/2 \\ -1/2 \\ 0 \end{bmatrix}}{\sqrt{1/4+1/4}} = \begin{bmatrix} 1/\sqrt{2} \\ -1/\sqrt{2} \\ 0 \end{bmatrix}$$

$$\begin{aligned} \underline{z}_3 &= \underline{x}_3 - (\underline{x}_3, \underline{e}_1) \underline{e}_1 - (\underline{x}_3, \underline{e}_2) \underline{e}_2 \\ &= \begin{bmatrix} 1/3 \\ 1/3 \\ -1/3 \end{bmatrix} \end{aligned}$$

$$\underline{e}_3 = \begin{bmatrix} 1/\sqrt{3} \\ 1/\sqrt{3} \\ -1/\sqrt{3} \end{bmatrix}$$

$$11) \quad y_1 = \frac{x_1}{\sqrt{\text{var}(x_1)}} = x_1$$

$$\begin{aligned} z_2 &= x_2 - (x_2, y_1) y_1 = x_2 - E(x_2 y_1) y_1 \\ &= x_2 - E(x_1 x_2) x_1 = x_2 - \rho x_1 \end{aligned}$$

$$y_2 = \frac{x_2 - \rho x_1}{\sqrt{E((x_2 - \rho x_1)^2)}} = \frac{x_2 - \rho x_1}{\sqrt{1 - 2\rho^2 + \rho^2}} = \frac{x_2 - \rho x_1}{\sqrt{1 - \rho^2}}$$

$$z_3 = x_3 - (x_3, y_1) y_1 - (x_3, y_2) y_2$$

$$= x_3 - E(x_3 x_1) x_1 - E\left(x_3 \frac{x_2 - \rho x_1}{\sqrt{1-\rho^2}}\right) \frac{x_2 - \rho x_1}{\sqrt{1-\rho^2}}$$

$$= x_3 - \rho^2 x_1 - \frac{\rho - \rho^3}{\sqrt{1-\rho^2}} \frac{x_2 - \rho x_1}{\sqrt{1-\rho^2}}$$

$$= x_3 - \rho^2 x_1 - \rho(x_2 - \rho x_1) = x_3 - \rho x_2$$

$$y_3 = \frac{x_3 - \rho x_2}{\sqrt{E((x_3 - \rho x_2)^2)}} = \frac{x_3 - \rho x_2}{\sqrt{1-\rho^2}}$$

$$\underline{y} = \underbrace{\begin{bmatrix} 1 & 0 & 0 \\ \frac{-\rho}{\sqrt{1-\rho^2}} & \frac{1}{\sqrt{1-\rho^2}} & 0 \\ 0 & \frac{-\rho}{\sqrt{1-\rho^2}} & \frac{1}{\sqrt{1-\rho^2}} \end{bmatrix}}_{\underline{A}} \underline{x}$$

$\underline{A}$ is lower triangular and will be in general.

$$\text{Since } \underline{C}_{yy} = \underline{I} \Rightarrow \underline{C}_{yy} = \underline{A} \underline{C}_{xx} \underline{A}^T$$

$$\underline{C}_{xx} = \underline{A}^{-1} \underline{A}^{T-1} \Rightarrow \underline{C}_{xx}^{-1} = \underline{A}^T \underline{A}$$

Thus, $\underline{A}$ is a whitening transformation and $\underline{A}^T \underline{A}$ provides a Cholesky decomposition of $\underline{C}_{xx}^{-1}$.

- 12) From (12.47) for a scalar parameter so that $\Omega(n-1)$ and $b(n)$ are scalars, as $\sigma_n^2 \rightarrow 0$

$$K[n] \rightarrow \frac{M[n-1] h[n]}{M[n-1] h^2[n]} = \frac{1}{h[n]}$$

$$\Rightarrow \hat{\theta}[n] = \hat{\theta}[n-1] + \frac{1}{h[n]} (x[n] - h[n] \hat{\theta}[n-1]) \\ = x[n] / h[n]$$

We discard all previous data since $x[n]$ is a perfect measurement (no noise). Also,

$$M[n] = (1 - K[n] h[n]) M[n-1] = 0$$

In vector case we do not obtain analogous result since we require p noiseless measurements to determine $\underline{\theta}$.

13) Here $h[n] = 1$, $\sigma_n^2 = \sigma^2$

$$\Rightarrow \hat{A}[n] = \hat{A}[n-1] + K[n] (x[n] - \hat{A}[n-1])$$

$$K[n] = \frac{M[n-1]}{\sigma^2 + M[n-1]}$$

$$M[n] = (1 - K[n]) M[n-1]$$

This is the same as the introductory example of Sect 12.6 except for the initialization since $\hat{A}[1] = E(A) = 0$

$$M[1] = E[(A - \hat{A}[1])^2] = E(A^2) = A_0^2/3$$

Solving for $K[n]$:

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

$$= \frac{\sigma^2 M[n-1]}{\sigma^2 + M[n-1]} = K[n] \sigma^2$$

$$\Rightarrow K[n] = \frac{K[n-1] \sigma^2}{\sigma^2 + K[n-1] \sigma^2} = \frac{K[n-1]}{1 + K[n-1]}$$

$$\frac{1}{K[n]} = \frac{1}{K[n-1]} + 1 \quad n \geq 1$$

$$K[0] = \frac{M[-1]}{\sigma^2 + M[-1]} = \frac{A_0^2/3}{\sigma^2 + A_0^2/3}$$

$$\frac{1}{K[0]} = 1 + \frac{\sigma^2}{A_0^2/3}$$

$$\frac{1}{K[n]} = \frac{1}{K[0]} + n = (n+1) + \frac{\sigma^2}{A_0^2/3}$$

$$K[n] = \frac{A_0^2/3}{(n+1) A_0^2/3 + \sigma^2}$$

$$1 - K[n] = \frac{n A_0^2/3 + \sigma^2}{(n+1) A_0^2/3 + \sigma^2}$$

$$\hat{A}[n] = (1 - K[n]) \hat{A}[n-1] + K[n] x[n]$$

$$= \frac{n A_0^2/3 + \sigma^2}{(n+1) A_0^2/3 + \sigma^2} \hat{A}[n-1] + \frac{A_0^2/3}{(n+1) A_0^2/3 + \sigma^2} x[n]$$

$$\text{where } \hat{A}[-1] = 0$$

$$\hat{A}[0] = \frac{A_0^2/3}{A_0^2/3 + \sigma^2} x[0]$$

$$\hat{A}[1] = \frac{A_0^2/3 + \sigma^2}{2 A_0^2/3 + \sigma^2} \cdot \frac{A_0^2/3}{A_0^2/3 + \sigma^2} x[0]$$

$$\begin{aligned}
& + \frac{A_0^2/3}{2 A_0^2/3 + \sigma^2} x(1) \\
& = \frac{A_0^2/3}{2 A_0^2/3 + \sigma^2} (x(0) + x(1)) \\
& = \frac{A_0^2/3}{A_0^2/3 + \sigma^2/2} \underbrace{\frac{1}{2} (x(0) + x(1))}_{\bar{x}}
\end{aligned}$$

etc.

14) To minimize $E[(x(n) - \hat{x}(n))^2]$ we use the orthogonality principle or

$$E[(x(n) - \hat{x}(n))x(n-l)] = 0 \quad \begin{array}{l} l = -M, \dots, M \\ l \neq 0 \end{array}$$

$$\begin{aligned}
r_{xx}(l) &= E \left[\sum_k a_k x(n-k) x(n-l) \right] \\
&= \sum_k a_k r_{xx}(l-k)
\end{aligned}$$

To show that $a_{-k} = a_k$:

Let $k' = -k$

$$\begin{aligned}
r_{xx}(l) &= \sum_{\substack{k'=-M \\ k' \neq 0}}^M a_{-k'} r_{xx}(l+k') \\
&\quad \begin{array}{l} l = -M, \dots, M \\ l \neq 0 \end{array}
\end{aligned}$$

Let $l' = -l$

$$\begin{aligned}
r_{xx}(l-l') &= \sum_{\substack{k'=-M \\ k' \neq 0}}^M a_{-k'} r_{xx}(-l'+k') \\
&\quad \begin{array}{l} l' = -M, \dots, M \\ l' \neq 0 \end{array}
\end{aligned}$$

$$r_{xx}(l-l') = \sum_{\substack{k'=-M \\ k' \neq 0}}^M a_{-k'} r_{xx}(l'-k')$$

$$\text{Since } r_{xx}(-k) = r_{xx}(k)$$

But these are the same set of equations for which there is a unique solution. Hence, $a_{-k} = a_k$. This must be true since the correlation of $x(n)$ with $x(n+k)$ is the same as that with $x(n-k)$, due to the even symmetry of the ACF.

$$15) \quad W = \frac{r_{ss}(0)}{r_{ss}(0) + r_{ww}(0)} = \eta / \eta + 1$$

$$M\hat{s} = \left(1 - \frac{\eta}{\eta + 1}\right) r_{ss}(0)$$

$$\begin{aligned} \text{But } \rho &= \frac{E(s(0)(s(0) + w(0)))}{\sqrt{r_{ss}(0)(r_{ss}(0) + r_{ww}(0))}} \\ &= \frac{r_{ss}(0)}{\sqrt{r_{ss}(0)(r_{ss}(0) + r_{ww}(0))}} \\ &= \frac{1}{\sqrt{1 + 1/\eta}} = \sqrt{\eta / \eta + 1} \end{aligned}$$

$$\Rightarrow W = \rho^2 \quad \text{or} \quad \hat{s}(0) = \rho^2 x(0)$$

$$M\hat{s} = (1 - \rho^2) r_{ss}(0)$$

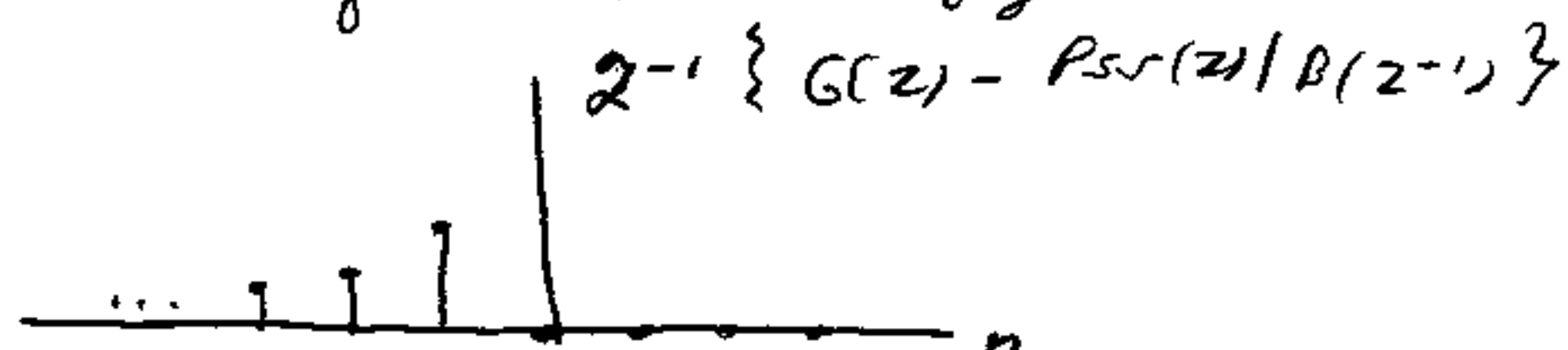
The higher the SNR, η , the larger is ρ and hence the better is the estimator.

$$16) \quad \left[B(z^{-1}) \left(G(z) - \frac{P_{sr}(z)}{B(z^{-1})} \right) \right]_+ = 0$$

But $B(z^{-1})$ is the z -transform of an anticausal sequence and thus if

$G(z) - \frac{P_{sr}(z)}{B(z^{-1})}$ is also anticausal

with a sequence value of zero at $n = 0$ or



the convolution will be zero for $n \geq 0$ as required. Now, since $P_{sr}(z)/B(z^{-1})$ is a two-sided sequence we let

$$z^{-1} \{ G(z) \} = z^{-1} \left\{ \frac{P_{sr}(z)}{B(z^{-1})} \right\} \text{ for } n \geq 0$$

$$\text{or } G(z) = \left[\frac{P_{sr}(z)}{B(z^{-1})} \right]_+$$

$$\text{or } H(z) = \frac{1}{B(z)} \left[P_{sr}(z)/B(z^{-1}) \right]_+$$

$$17) \quad E \{ (S(n) - \hat{S}(n)) x(n-l) \} = 0 \quad -\infty < l < \infty$$

$$E \{ (S(n) - \sum_k h[k] x(n-k)) x(n-l) \} = 0$$

$$E \{ S(n) x(n-l) \} = \sum_k h[k] E \{ x(n-k) x(n-l) \}$$

$$r_{sr}[l] = \sum_{k=-\infty}^{\infty} h[k] r_{xx}[l-k]$$

$$\begin{aligned}
 18) \quad M_{\hat{s}} &= E \{ (s[n] - \hat{s}[n])^2 \} \\
 &= E \{ (s[n] - \hat{s}[n]) (s[n] - \sum_k h[k] x[n-k]) \} \\
 &= E \{ (s[n] - \hat{s}[n]) s[n] \} \\
 &\quad - E \{ (s[n] - \hat{s}[n]) \sum_k h[k] x[n-k] \} \\
 &\qquad\qquad\qquad = 0 \text{ by orthogonality principle} \\
 &= r_{ss}[0] - \sum_k h[k] \underbrace{E \{ x[n-k] s[n] \}}_{r_{sx}[k]}
 \end{aligned}$$

Now by Parseval's theorem

$$\sum_k h[k] r_{sx}[k] = \int P_{ss}(f) H(f) df$$

$$\begin{aligned}
 \Rightarrow M_{\hat{s}} &= \int P_{ss}(f) df - \int P_{ss}(f) H(f) df \\
 &= \int_{-\frac{1}{2}}^{\frac{1}{2}} P_{ss}(f) (1 - H(f)) df
 \end{aligned}$$

For the example given

$$\begin{aligned}
 H(f) &= \frac{P_0}{P_0 + \sigma^2} & |f| \leq 1/4 \\
 &0 & 1/4 < |f| \leq 1/2
 \end{aligned}$$

$$M_{\hat{s}} = \int_{-1/4}^{1/4} P_0 \left(1 - \frac{P_0}{P_0 + \sigma^2} \right) df = \frac{1}{2} \frac{P_0 \sigma^2}{P_0 + \sigma^2}$$

Since there is no signal power above $f = 1/4$, $H(f) = 0$. For $f \leq 1/4$ we weight all frequencies (since signal and noise have flat PSDs) by an SNR weighting or

$$H(f) = \frac{P_0}{P_0 + \sigma^2} = \frac{\eta}{\eta + 1}$$

$$19) \quad \hat{x}(n) = \sum_{k=1}^N h(k) x(n-k)$$

$$E[(x(n) - \hat{x}(n)) x(n-l)] = 0 \quad l = 1, 2, \dots, N$$

$$\begin{aligned} r_{xx}(l) &= \sum_{k=1}^N h(k) E(x(n-k) x(n-l)) \\ &= \sum_{k=1}^N h(k) r_{xx}(l-k) \end{aligned}$$

The equations are independent of n (in deriving (12.65) we assumed $n=N$ was the index of the sample to be predicted) since the ACF does not depend on n .

$$\begin{aligned} M_2 &= E[(x(n) - \hat{x}(n)) x(n)] \\ &\quad - E[(x(n) - \hat{x}(n)) \hat{x}(n)] \\ &= 0 \text{ by orthogonality principle} \end{aligned}$$

$$= E[x^2(n)] - \sum_{k=1}^N h(k) E[(n-k) x(n)]$$

$$= r_{xx}(0) - \sum_{k=1}^N h(k) r_{xx}(k)$$

20) From the previous problem we have

$$r_{xx}(l) = \sum_{k=1}^N h(k) r_{xx}(l-k) \quad l = 1, 2, \dots, N$$

must be solved for the optimal one-step predictor. But for an AR(N) process we know that (see Appendix 1)

$$r_{xx}(k) = - \sum_{h=1}^N a[h] r_{xx}(k-h) \quad k \geq 1$$

which are the Yule-Walker equations. Since the solution for the $h[k]$'s is unique,

$$h[k] = -a[k]$$

so that $\hat{x}[n] = - \sum_{k=1}^N a[k] x[n-k]$ and the MMSE is

$$M_x^{\wedge} = r_{xx}[0] - \sum_{k=1}^N h[k] r_{xx}[k]$$

$$= r_{xx}[0] + \sum_{k=1}^N a[k] r_{xx}[k]$$

$$= \sigma_u^2 \quad (\text{see Appendix 1})$$

Chapter 13

1)
$$s[n] = a^{n+1} s[-1] + \sum_{k=0}^n a^k u[n-k]$$

Let $\underline{s} = [s[n_1] \ s[n_2] \ \dots \ s[n_k]]^T$

Assume $n_k > n_{k-1} > \dots > n_1$

$$\underline{s} = \begin{bmatrix} a^{n_1+1} & a^{n_1} & \dots & 1 & 0 & 0 & \dots & 0 \\ a^{n_2+1} & a^{n_2} & \dots & & 1 & 0 & \dots & 0 \\ \vdots & & & & & & & \\ a^{n_k+1} & a^{n_k} & \dots & & & & & 1 \end{bmatrix} \begin{bmatrix} s[-1] \\ u[0] \\ \vdots \\ u[n_k] \end{bmatrix}$$

Since $s[-1], u[0], \dots, u[n_k]$ are independent, they are jointly Gaussian and thus $\underline{s}$ has a multivariate Gaussian PDF, being a linear transformation.

2) If $u_s = 0$, from (13.4), $E(s[n]) = 0$

$$\begin{aligned} E[s[n]] &= a^{n+1} \sigma_s^2 + \sigma_u^2 a^{n+1} \sum_{k=0}^n a^{2k} \\ &= \frac{a^{n+1} \sigma_u^2}{1-a^2} + \frac{\sigma_u^2 a^{n+1} (1-a^{2(n+1)})}{1-a^2} \\ &= \frac{\sigma_u^2}{1-a^2} a^{n+1} \end{aligned}$$

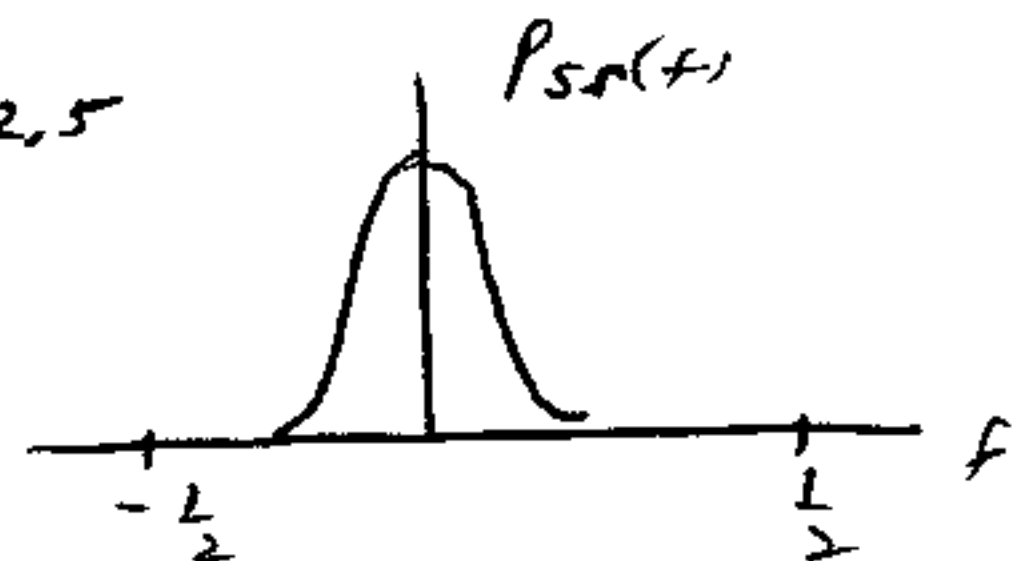
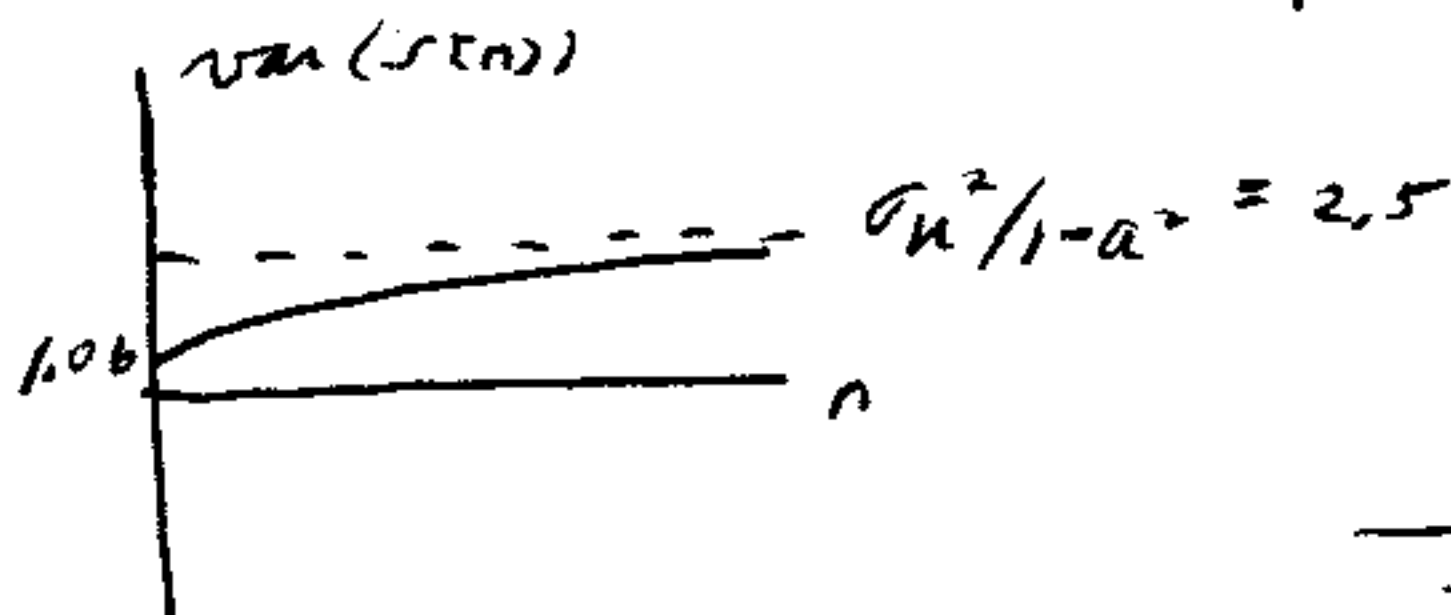
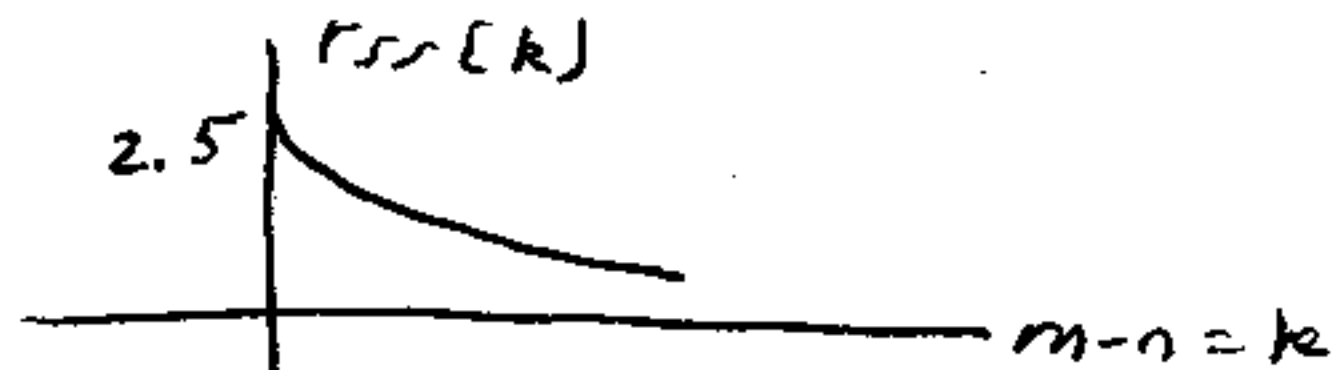
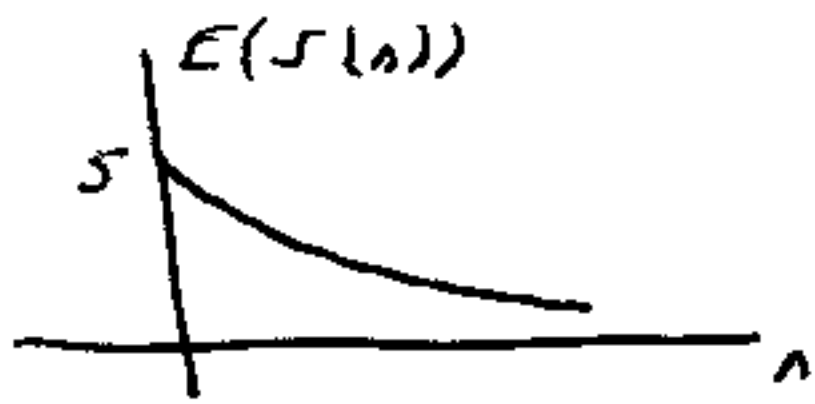
$\Rightarrow$ WSS. By setting the initial conditions as given the process is in steady-state for $n \geq -1$.

$$\begin{aligned}
 3) \quad E(S(n)) &= a^{n+1} \mu_s = 5(0.98)^{n+1} \\
 \text{var}(S(n)) &= a^{2n+2} \sigma_s^2 + \sigma_u^2 \sum_{k=0}^n a^{2k} \\
 &= (0.98)^{2n+2} + 0.1 \sum_{k=0}^n (0.98)^{2k}
 \end{aligned}$$

$$\begin{aligned}
 C_S(m, n) &\rightarrow \frac{\sigma_u^2}{1-a^2} a^{m-n} \\
 &= \frac{0.1}{1-0.98^2} 0.98^{m-n}
 \end{aligned}$$

PSD is just that of an AR(1) process or

$$\begin{aligned}
 P_{SS}(f) &= \frac{\sigma_u^2}{|1 - a e^{-j2\pi f}|^2} \\
 &= \frac{0.1}{|1 - 0.98 e^{-j2\pi f}|^2}
 \end{aligned}$$



4) From (13.5)

$$C_S(m, n) = a^{m-n} \underbrace{\left[a^{2n+2} \sigma_s^2 + \sigma_u^2 \sum_{k=0}^n a^{2k} \right]}_{\text{var}(S(n))}$$

$$C_S(m, n) = a^{m-n} C_S(n, n)$$

$$5) \quad \underline{z}[n] = \underline{A} \underline{z}[n-1] + \underline{B} \underline{u}[n]$$

$$\begin{aligned} E(\underline{z}[n]) &= \underline{A} E(\underline{z}[n-1]) + \underline{B} E(\underline{u}[n]) \\ &= \underline{A} E(\underline{z}[n-1]) \end{aligned}$$

$$\begin{aligned} C[n] &= E[(\underline{z}[n] - E(\underline{z}[n]))(\underline{z}[n] - E(\underline{z}[n]))^T] \\ &= E[(\underline{A} \underline{z}[n-1] + \underline{B} \underline{u}[n] - \underline{A} \overset{E(\underline{z}[n-1])}{\underline{z}[n-1]}) \\ &\quad (\underline{A} \underline{z}[n-1] + \underline{B} \underline{u}[n] - \underline{A} \overset{E(\underline{z}[n-1])}{\underline{z}[n-1]})^T] \end{aligned}$$

$$= E[(\underline{A} (\underline{z}[n-1] - E(\underline{z}[n-1])) + \underline{B} \underline{u}[n]) \\ (\underline{A} (\underline{z}[n-1] - E(\underline{z}[n-1])) + \underline{B} \underline{u}[n])^T]$$

$$= \underline{A} C[n-1] \underline{A}^T + \underline{B} \underline{Q} \underline{B}^T \quad \text{since}$$

$$E(\underline{z}[n-1] \underline{u}[n]^T) = \underline{0}$$

($\underline{z}[n-1]$ depends on $\underline{u}[k]$ for $k \leq n-1$).

$$6) \quad \text{From (3.14)} \quad E(\underline{z}[n]) = \underline{A}^{n+1} \underline{\mu}_s$$

Let $\underline{V}$ be the modal matrix for $\underline{A}$

and $\underline{\Lambda}$ the diagonal matrix of eigenvalues.

$$\text{Thus, } \underline{V}^T \underline{A} \underline{V} = \underline{\Lambda} \quad \text{or} \quad \underline{A} = \underline{V} \underline{\Lambda} \underline{V}^T$$

$$\Rightarrow \underline{A}^{n+1} = \underline{V} \underline{\Lambda}^{n+1} \underline{V}^T$$

$$\begin{aligned} E(\underline{z}[n]) &= \underline{V} \underline{\Lambda}^{n+1} \underline{V}^T \underline{\mu}_s \\ &= \sum_{i=1}^p a_i \lambda_i^{n+1} \underline{v}_i \end{aligned}$$

where $a_i = (\underline{v}^T \underline{u}_i)_i$
 $\underline{v} = (\underline{v}_1 \underline{v}_2 \dots \underline{v}_p)$

If any $|\lambda_i| > 1 \quad E(\underline{\xi}(n)) \rightarrow \infty$

If all $|\lambda_i| < 1 \quad E(\underline{\xi}(n)) \rightarrow 0$

7) $\underline{A}^n = (\underline{V} \underline{\Lambda} \underline{V}^T)^n = \underline{V} \underline{\Lambda}^n \underline{V}^T$
 $\underline{A}^n \underline{e}_i \underline{A}^{nT} = \underline{V} \underline{\Lambda}^n \underline{V}^T \underline{e}_i \underline{V} \underline{\Lambda}^n \underline{V}^T$
 $\underline{e}_i^T \underline{A}^n \underline{e}_j \underline{A}^{nT} \underline{e}_j = \underbrace{\underline{e}_i^T \underline{V} \underline{\Lambda}^n (\underline{V}^T \underline{e}_j \underline{V})}_{\underline{b}^T} \underbrace{\underline{\Lambda}^n \underline{V}^T \underline{e}_j}_{\underline{a}}$

But $\underline{b} = \underline{\Lambda}^n \underline{V} \underline{e}_i \rightarrow 0$ if $|\lambda_i| < 1$

$\underline{a} = \underline{\Lambda}^n \underline{V} \underline{e}_j \rightarrow 0$ if $|\lambda_i| < 1$

$\Rightarrow \underline{b}^T \underline{V}^T \underline{e}_i \underline{V} \underline{a} = [\underline{A}^n \underline{e}_i \underline{A}^{nT}]_{ij} \rightarrow 0$
for all i, j

8)
$$\underbrace{\begin{bmatrix} r(n-p+1) \\ r(n-p+2) \\ \vdots \\ r(n) \end{bmatrix}}_{\underline{\xi}(n)} = \underbrace{\begin{bmatrix} 0 & 1 & 0 & \dots & 0 \\ 0 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ -a(p) & -a(p-1) & \dots & -a(1) \end{bmatrix}}_{\underline{A}} \underbrace{\begin{bmatrix} r(n-p) \\ r(n-p+1) \\ \vdots \\ r(n-1) \end{bmatrix}}_{\underline{\xi}(n-1)}$$

$$+ \underbrace{\begin{bmatrix} 0 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ b(q) & b(q-1) & \dots & 1 \end{bmatrix}}_{\underline{B}} \underbrace{\begin{bmatrix} u(n-q) \\ u(n-q+1) \\ \vdots \\ u(n) \end{bmatrix}}_{\underline{u}(n)}$$

Not a Gauss-Markov process since $\underline{u}(n)$ is not vector WGN. This is because $\underline{u}(n)$ is correlated in time. If $q=1$ for example

$$\begin{aligned} E(\underline{u}(n) \underline{u}(n+1)^T) &= E \begin{bmatrix} u(n-1) \\ u(n) \end{bmatrix} \begin{bmatrix} u(n) & u(n+1) \end{bmatrix} \\ &= \begin{bmatrix} E[u(n-1)u(n)] & E[u(n-1)u(n+1)] \\ E[u^2(n)] & E[u(n)u(n+1)] \end{bmatrix} \\ &= \begin{bmatrix} 0 & 0 \\ \sigma_u^2 & 0 \end{bmatrix} \neq 0 \end{aligned}$$

$$9) \quad r(n) = s(n) + \sum_{l=1}^q b(l) s(n-l)$$

$$\begin{aligned} \sum_{k=0}^p a(k) r(n-k) &= \sum_{k=0}^p a(k) s(n-k) \\ &\quad + \sum_{l=1}^q b(l) \sum_{k=0}^p a(k) s(n-k+l) \end{aligned}$$

$$= u(n) + \sum_{l=1}^q b(l) u(n-l)$$

Now let the state

be given by (13.10)

$$\Rightarrow \underline{s}(n) = \underline{A} \underline{s}(n-1) + \underline{B} u(n) \quad (\text{see development leading to (13.11)})$$

$$\text{Since } \underline{s}(n) = \begin{bmatrix} s(n-p+1) \\ s(n-p+2) \\ \vdots \\ s(n) \end{bmatrix}, \quad \text{if } q \leq p-1$$

$$r(n) = \underbrace{(0 \dots 0 \ b(q) \dots b(1) \ 1)}_{h^T(n) = h^T} \underline{s}(n)$$

$$\Rightarrow x(n) = \underline{h}^T \underline{s}(n) + \underline{w}(n)$$

Now use (13.51) - (13.56) with

$$\underline{A} = \begin{bmatrix} 0 & 1 & 0 & \dots & 0 \\ 0 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ -a(p) & -a(p-1) & \dots & -a(1) & 0 \end{bmatrix} \quad \underline{B} = \begin{bmatrix} 0 \\ \vdots \\ 0 \\ 1 \end{bmatrix}$$

$$\underline{Q} = \sigma_w^2$$

10) From (13.38) - (13.42) with $a=1$, $\sigma_w^2=0$ so that $s(n) = s(n-1) = A$. We can skip the prediction stage since

$$\hat{A}(n|n-1) = \hat{A}(n-1|n-1)$$

$$M(n|n-1) = M(n-1|n-1)$$

$$\Rightarrow K(n) = \frac{M(n-1|n-1)}{\sigma^2 + M(n-1|n-1)}$$

$$\hat{A}(n|n) = \hat{A}(n-1|n-1) + K(n)(x(n) - \hat{A}(n-1|n-1))$$

$$M(n|n) = (1 - K(n))M(n-1|n-1)$$

or changing the notation we have

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

$$\hat{A}(n) = \hat{A}(n-1) + K(n)(x(n) - \hat{A}(n-1))$$

$$M(n) = (1 - K(n))M(n-1)$$

These equations are just (12.34) - (12.36) with obvious changes in notation.

Hence, from Section 12.6

$$\hat{A}(n) = \frac{\sigma_A^2}{\sigma_A^2 + \sigma^2/n+1} \quad \frac{1}{n+1} \sum_{k=0}^n x[k]$$

$$M(n-1) = \frac{\sigma_A^2 \sigma^2}{n \sigma_A^2 + \sigma^2}$$

$$K(n) = \frac{\sigma_A^2}{(n+1) \sigma_A^2 + \sigma^2}$$

11) From (13.39), (13.40), (13.42)

$$M(n|n-1) = 0.81 M(n-1|n-1) + 1$$

$$K(n) = \frac{M(n|n-1)}{\sigma_n^2 + M(n|n-1)}$$

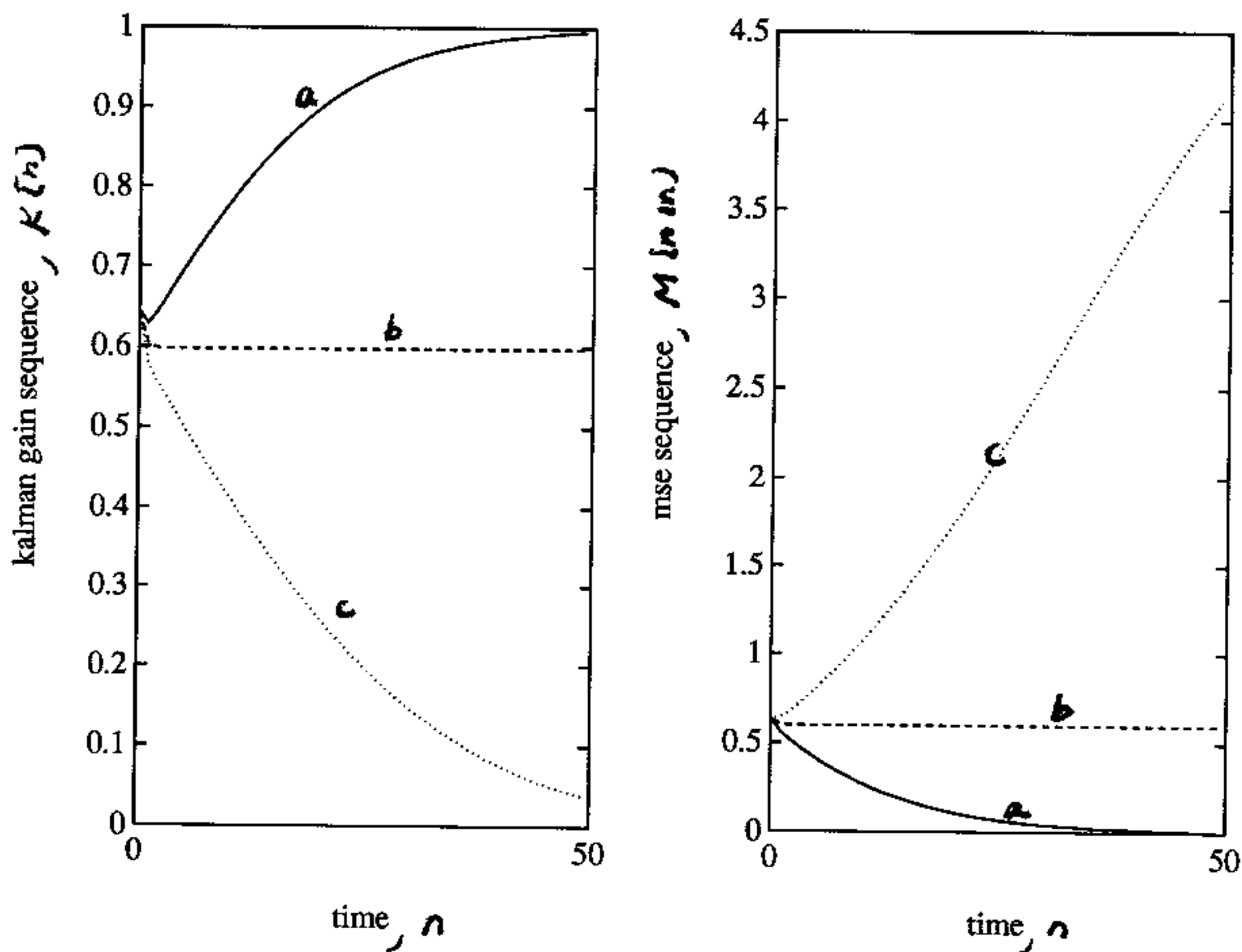
$$M(n|n) = (1 - K(n)) M(n|n-1)$$

$$\text{where } M(-1|-1) = \sigma_s^2 = 1$$

See next page for plots. For (a) the gain $\rightarrow 1$ since the observations become less noisy and thus $\hat{s}(n|n) \rightarrow x(n)$. Also, the signal estimate improves with time since $M(n|n) \rightarrow 0$. For (c) the gain $\rightarrow 0$ since the observations become noisier with time, and also $\hat{s}(n|n) \rightarrow \hat{s}(n|n-1) = 0.9 \hat{s}(n-1|n-1)$.

The signal estimate will decay to zero so that the MMSE will be $E[\hat{s}^2(n)] = \sigma_u^2 / 1 - a^2 = 5.26$, which is the steady-state variance of $s(n)$. For (b) the gain and MMSE sequences attain a steady state after a few

iterations. This is the steady-state Kalman filter or Wiener filter (see Section 13.5).

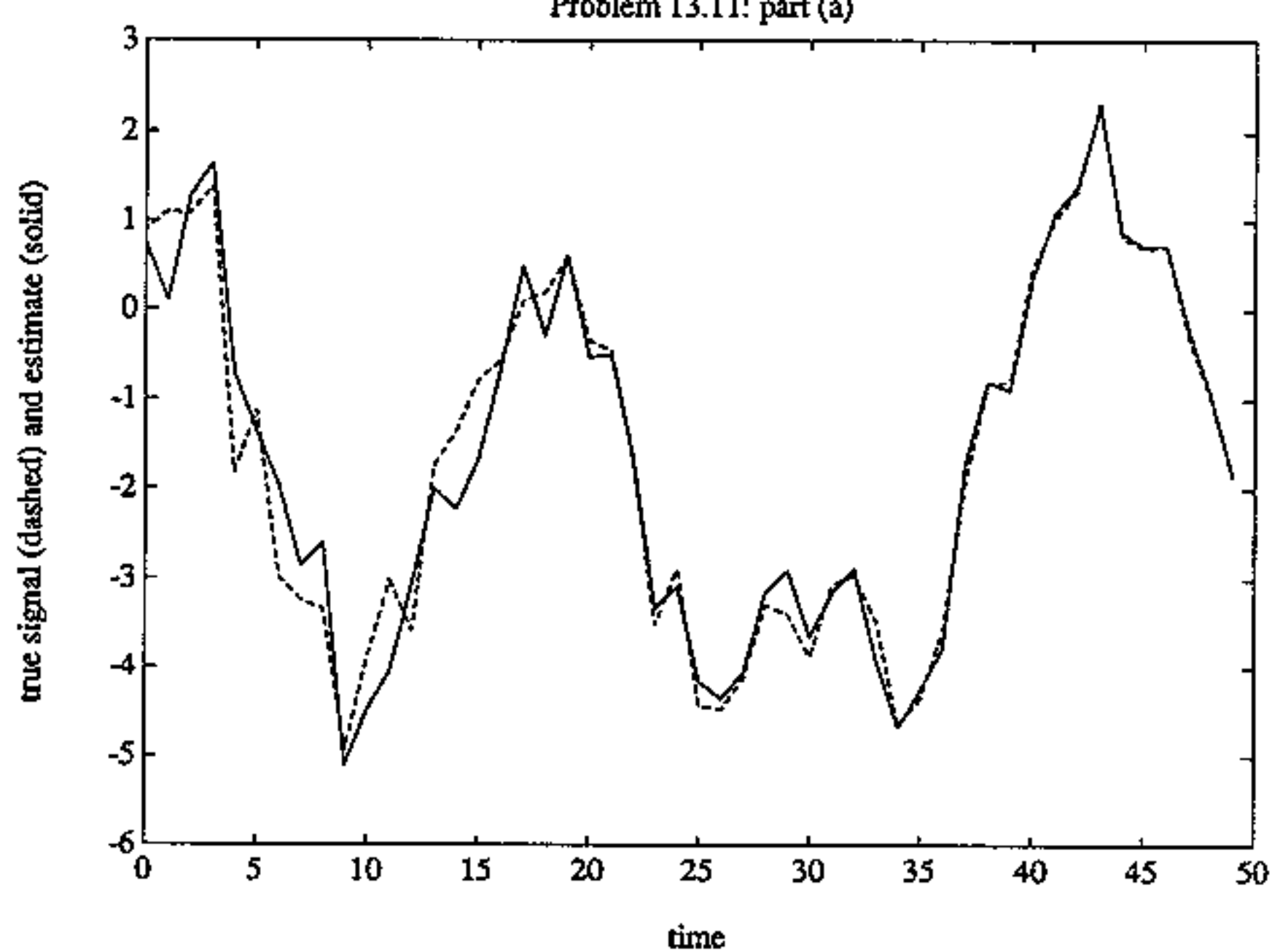


A Monte Carlo simulation produces the plots on the following page.

12) If $\sigma_n^2 = 0$, $K[n] = 1$ and $\hat{z}[n|n] = x[n]$
 Also, $\hat{z}[n|n-1] = a \hat{z}[n-1|n-1]$
 $= a x[n-1] = a \hat{z}[n-1]$
 Thus, $\tilde{z}[n] = x[n] - \hat{z}[n|n-1]$

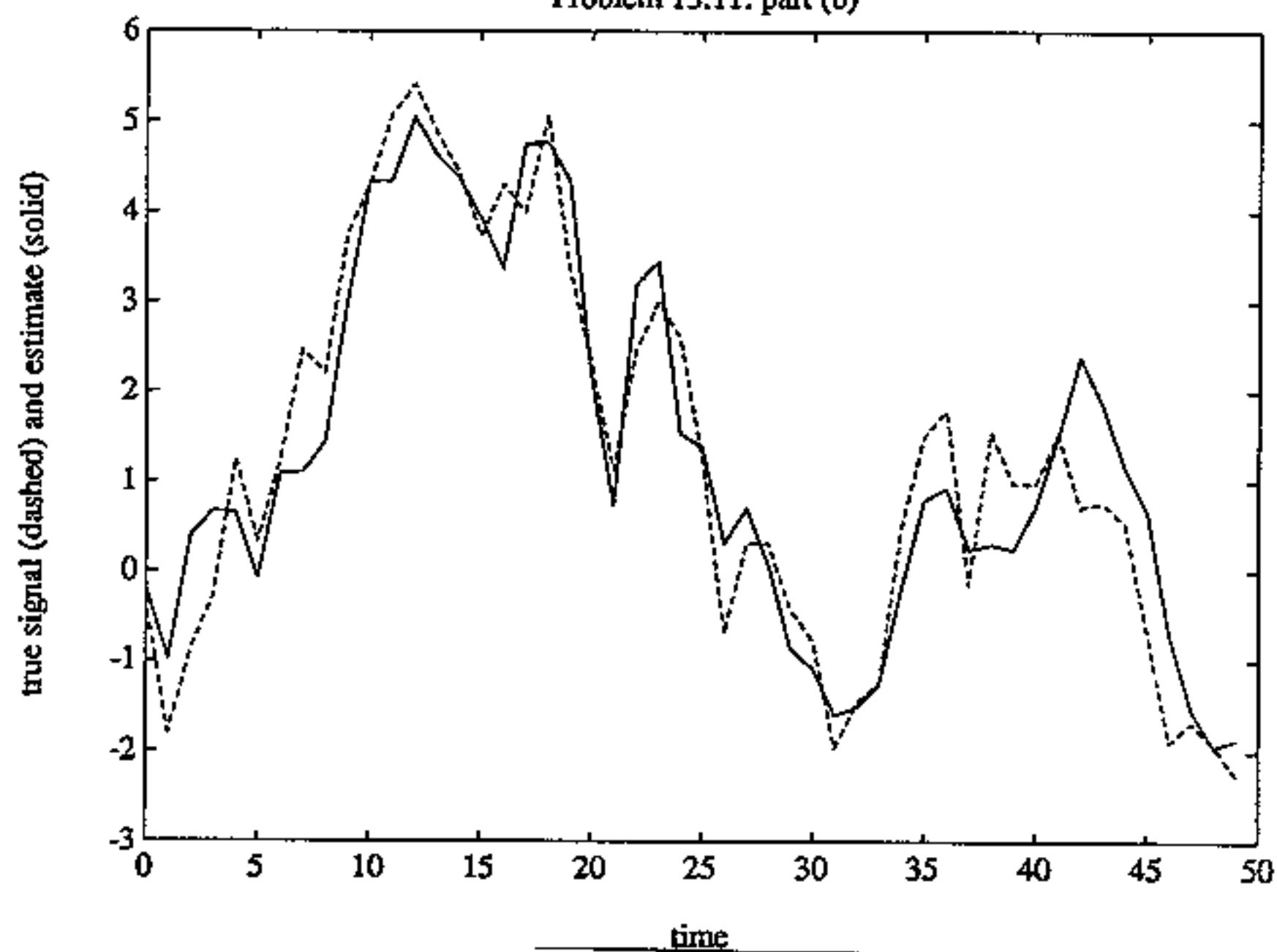
Problem 13.11: part (a)

191

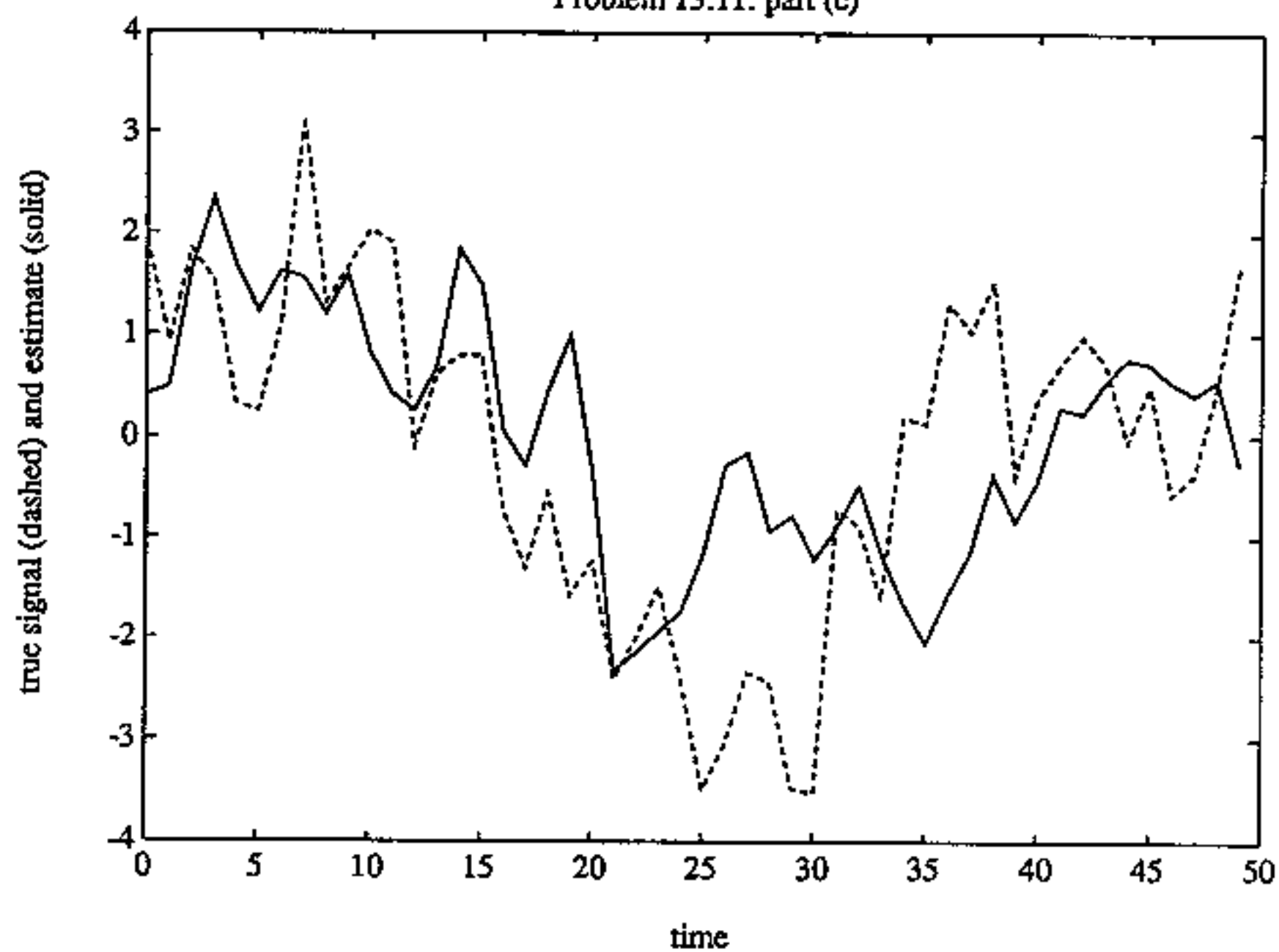


$s(n) = \text{DASHED}$
 $\hat{s}(n|n) = \text{SOLID}$

Problem 13.11: part (b)



Problem 13.11: part (c)



$$= s[n] - a s[n-1] = u[n]$$

Yes.

- 13) Since $E(s[n])$ is known, the MMSE estimator of $s'[n] = s[n] - E(s[n])$ is $\hat{s}'[n] = \hat{s}[n] - E(s[n])$ and the minimum MSE is the same as for $\hat{s}[n]$. Thus, $M[n|n]$ and $M[n|n-1]$ do not change. Also, then $K[n]$ does not change.

Prediction: $\hat{s}'[n|n-1] = a \hat{s}'[n-1|n-1]$

$$\hat{s}[n|n-1] - E(s[n]) = a (\hat{s}[n-1|n-1] - E(s[n-1]))$$

$$\text{or } \hat{s}[n|n-1] = a \hat{s}[n-1|n-1]$$

$$\text{Since } E(s[n]) = a E(s[n-1])$$

Correction: $\hat{s}'[n|n] = \hat{s}'[n|n-1]$

$$+ K[n] (x'[n] - \hat{s}'[n|n-1])$$

$$\hat{s}[n|n] - E(s[n]) = \hat{s}[n|n-1] - E(s[n])$$

$$+ K[n] (x[n] - E(x[n]) - \hat{s}[n|n-1] + E(s[n]))$$

$$\text{But } E(x[n]) = E(s[n]) + E(w[n]) = E(s[n])$$

Thus, we have the same equations as before.

The only difference arises in the initialization

since $u_s \neq 0$.

- 14) In steady state we have $M[n|n] = M[\infty]$, $M[n|n-1] = M_p[\infty]$ and from (13.42)

$$M[\infty] = (1 - K[\infty]) M_p[\infty] \\ < M_p[\infty]$$

since $K[\infty] < 1$. Thus, for large n
 $M[n|n-1] > M[n-1|n-1]$. This is
 reasonable since $s[n]$ is harder to
 estimate than $s[n-1]$ based on $\{x[0], x[1], \dots, x[n-1]\}$ due to the added variability
 of the $u[n]$ noise term.

- 15) From (13.38) $\hat{s}[n+1|n] = a \hat{s}[n|n]$
 Now if $\sigma_{n+1}^2 \rightarrow \infty$, the future measurements will be useless so that the
 corrected estimates will be predictions or
 will be based on only $\{x[0], \dots, x[n]\}$.

$$\Rightarrow \hat{s}[n+1|n+1] \rightarrow \hat{s}[n+1|n] = a \hat{s}[n|n] \\ \hat{s}[n+2|n+1] = a \hat{s}[n+1|n+1] = a^2 \hat{s}[n|n] \\ \text{etc.}$$

- 16) From (13.47)

$$H_{\infty}(z) = \frac{K[\infty]}{1 - a(1 - K[\infty])z^{-1}}$$

From (13.46)

$$M[\infty] = \frac{0.64 M[\infty] + 1}{0.64 M[\infty] + 2}$$

Let $x = M(\infty)$

$$0.64x^2 + 2x = 0.64x + 1$$

$$x^2 + 2.125x - 1.5625 = 0$$

$$\Rightarrow x = 0.5781, -2.703$$

$$M(\infty) = 0.5781$$

$$M_p(\infty) = a^2 M(\infty) + \sigma_u^2$$

$$= 0.64 M(\infty) + 1 = 1.37$$

$$K(\infty) = \frac{M_p(\infty)}{1 + M_p(\infty)} = 0.5781$$

$$H_{\infty}(z) = \frac{0.5781}{1 - 0.3375z^{-1}}$$

$$\hat{f}(n|n) = 0.3375 \hat{f}(n-1|n-1) + 0.5781 x[n]$$

17) See plot on next page.

18) Here $h(n) = r^n$. Using (13.51) - (13.54)

$$x(n) = h(n)A + w(n)$$

$$\underline{Q} = 0$$

$$\underline{A} = \underline{I}$$

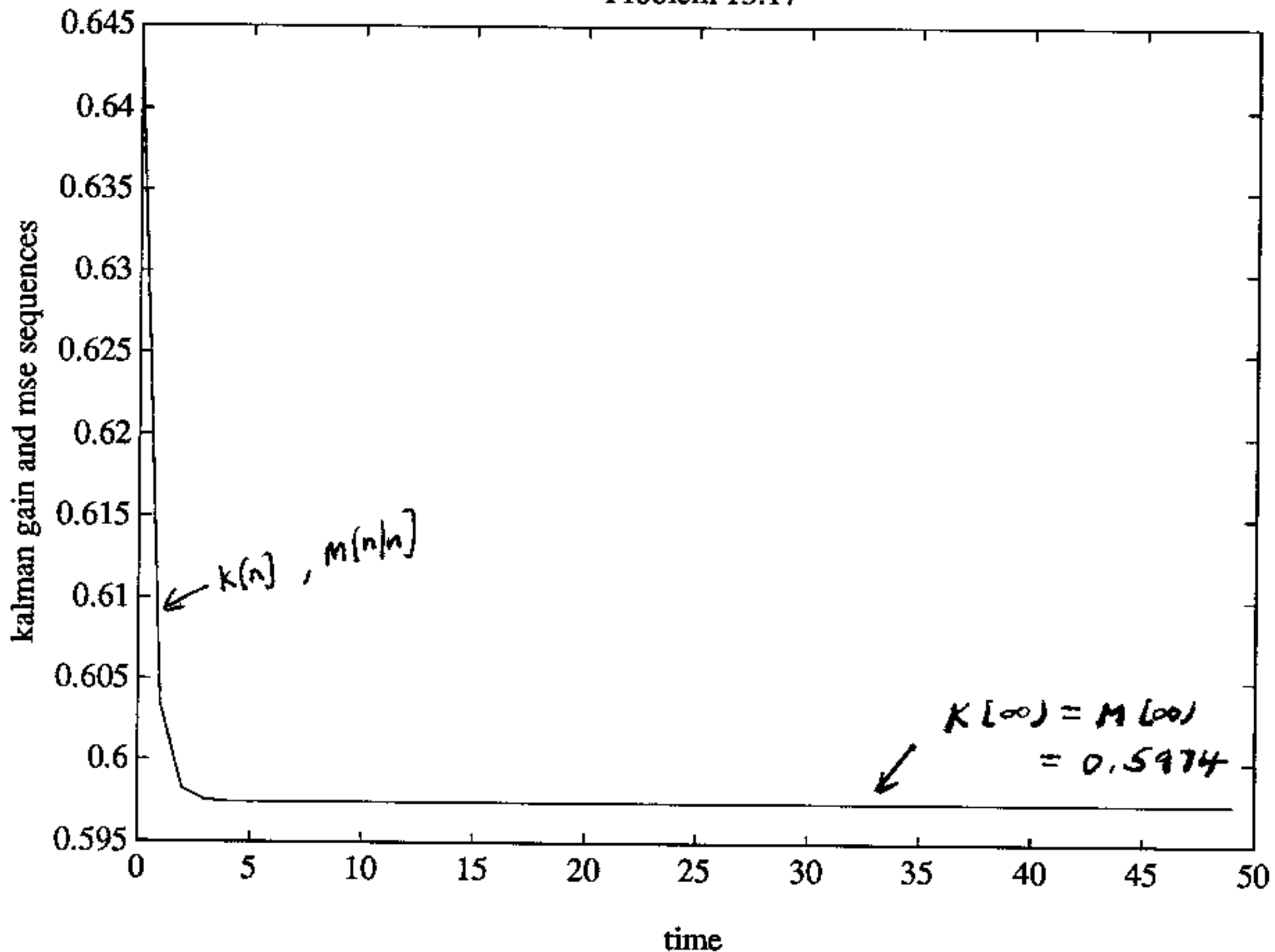
Can omit prediction stage.

$$K(n) = \frac{M(n|n-1)r^n}{\sigma^2 + r^{2n}M(n|n-1)}$$

$$M(n|n) = (1 - K(n)r^n) M(n|n-1)$$

$$\hat{A}(n|n) = \hat{A}(n|n-1) + K(n)(x(n) - r^n \hat{A}(n|n-1))$$

Problem 13.17



where $\hat{A}[-1|-1] = \mu_A$, $M[-1|-1] = \sigma_A^2$.

19) If $\underline{x}(n) = \underline{0}$, from (13.60)

$$\begin{aligned} \underline{K}(n) &= \underline{M}(n|n-1) \underline{H}^T(n) (\underline{H}(n) \underline{M}(n|n-1) \underline{H}^T(n) + \sigma_v^2)^{-1} \\ &= \underline{M}(n|n-1) \underline{H}^T(n) \underline{H}^{-1}(n) \underline{M}^{-1}(n|n-1) \underline{H}^{-1}(n) \\ &= \underline{H}^{-1}(n) \end{aligned}$$

$$\begin{aligned} \hat{\underline{x}}(n|n) &= \hat{\underline{x}}(n|n-1) + \underline{H}^{-1}(n) (\underline{x}(n) - \underline{H}(n) \hat{\underline{x}}(n|n-1)) \\ &= \underline{H}^{-1}(n) \underline{x}(n) \end{aligned}$$

$\Rightarrow$ discard all previous data since

$$\underline{x}(n) = \underline{H}(n) \underline{z}(n) \text{ and thus } \underline{z}(n) = \underline{H}^{-1}(n) \underline{x}(n)$$

$\Rightarrow$ no error in $\hat{\underline{x}}(n|n)$.

If $L(n) \rightarrow \infty$, $K(n) \rightarrow 0 \Rightarrow \hat{f}_0[n|n] \rightarrow \hat{f}_0[n|n-1]$
or we ignore the data sample $x[n]$.

20) Same approach as in Problem 13.15.

21) Note that the state equation is linear but the observation equation is not.

From (13.67) - (13.71)

$$\hat{f}_0[n|n-1] = a \hat{f}_0[n-1|n-1]$$

$$M[n|n-1] = a^2 M[n-1|n-1] + \sigma_u^2$$

$$K[n] = \frac{M[n|n-1] H[n]}{\sigma^2 + H^2[n] M[n|n-1]}$$

where $H[n] \triangleq \left. \frac{\partial h}{\partial f_0[n]} \right|_{f_0[n] = \hat{f}_0[n|n-1]}$

But $h(f_0[n]) = \cos 2\pi f_0[n]$

$$H[n] = -2\pi \sin 2\pi \hat{f}_0[n|n-1]$$

$$\hat{f}_0[n|n] = \hat{f}_0[n|n-1] + K[n] (x[n] - \cos 2\pi \hat{f}_0[n|n-1])$$

$$M[n|n] = (1 - K[n] H[n]) M[n|n-1]$$

22) $Nx[n] = Nx[n-1] + u_x[n]$

$$\Rightarrow Nx[n] = \sum_{k=0}^n u_x[k] + Nx[-1]$$

$$\text{var}(Nx[n]) = \sum_{k=0}^n \text{var}(u_x[k]) + \text{var}(Nx[-1])$$

Since $u_x[n]$ is WGN and is independent of $Nx[-1]$

$$\begin{aligned}\text{var}(x_n) &= (n+1) \text{var}(u_n) + \sigma_v^2 \\ &= (n+1) \sigma_u^2 + \sigma_v^2\end{aligned}$$

A better model would be

$$x_n = a x_{n-1} + u_n$$

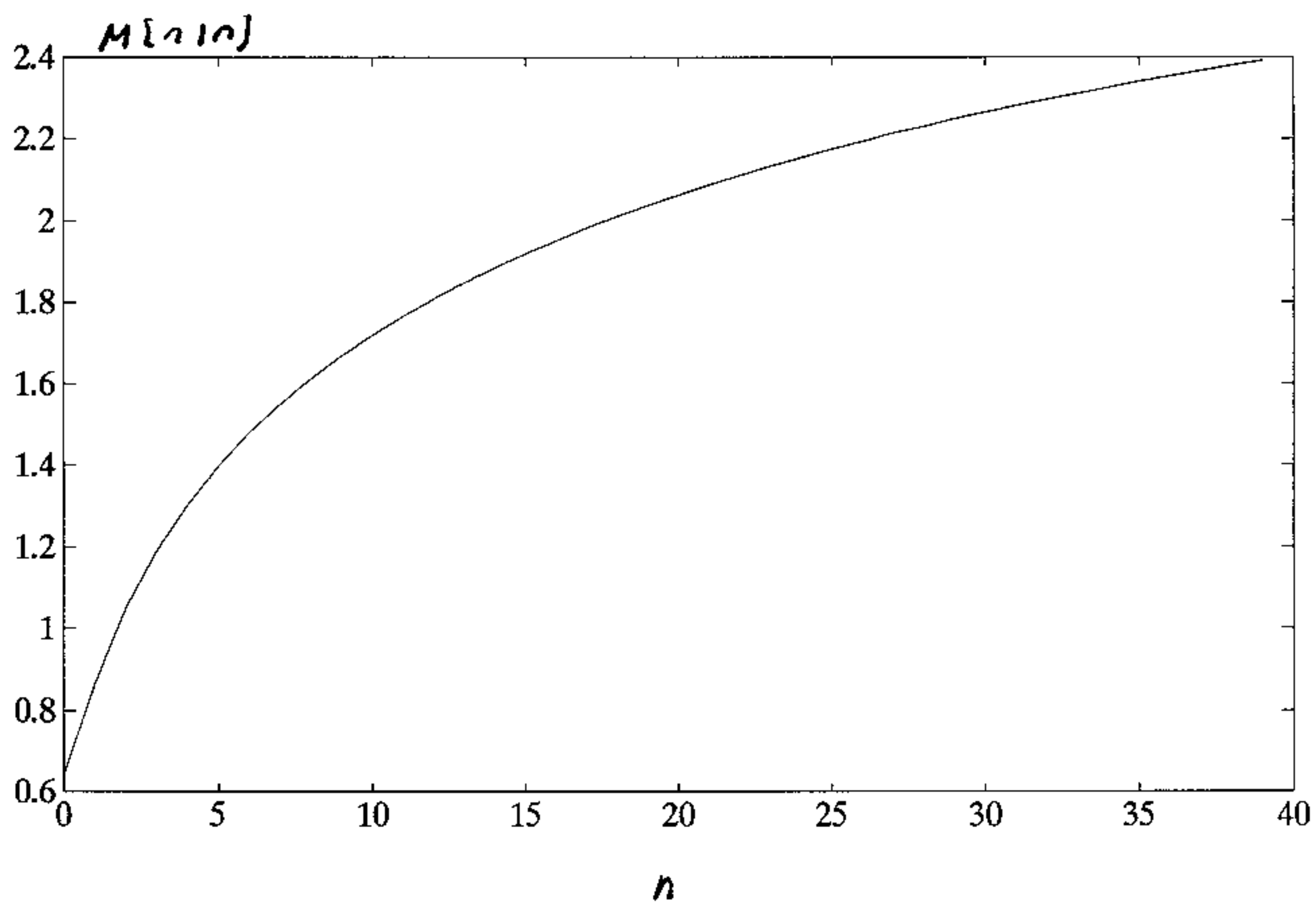
for $0 < a < 1$ and a should be near one.

Then, in steady-state the variance would be a constant.

$$\begin{aligned}23) \quad M[n|n] &= (1 - K[n]) M[n|n-1] \\ &= \left(1 - \frac{M[n|n-1]}{\sigma_n^2 + M[n|n-1]} \right) M[n|n-1] \\ &= \frac{\sigma_n^2 M[n|n-1]}{\sigma_n^2 + M[n|n-1]} \\ &= \frac{\sigma_n^2 (a^2 M[n-1|n-1] + \sigma_u^2)}{\sigma_n^2 + a^2 M[n-1|n-1] + \sigma_u^2} \\ &= \frac{(n+1) (0.81 M[n-1|n-1] + 1)}{n+1 + 0.81 M[n-1|n-1] + 1}\end{aligned}$$

where $M[-1|-1] = 1$

See plot on next page. Since the observations are progressively noisier, the MMSE increases.



Chapter 15

$$1) \quad \text{Cov}(\tilde{x}_1, \tilde{x}_2) = E((\tilde{x}_1 - E(\tilde{x}_1))^* (\tilde{x}_2 - E(\tilde{x}_2)))$$

The expectation is with respect to $p(u_1, v_1, u_2, v_2)$

$$\text{But } p(u_1, v_1, u_2, v_2) = p(u_1, v_1) p(u_2, v_2)$$

Thus,

$$\begin{aligned} \text{Cov}(\tilde{x}_1, \tilde{x}_2) &= E_{u_1, v_1}(\tilde{x}_1 - E(\tilde{x}_1)) E_{u_2, v_2}(\tilde{x}_2 - E(\tilde{x}_2)) \\ &= 0 \end{aligned}$$

$$2) \quad \text{Consider } \tilde{y} = \underline{a}^H \tilde{x}$$

$$E(|\tilde{y}|^2) = E(\underline{a}^H \tilde{x} \tilde{x}^H \underline{a}) = \underline{a}^H \underline{C}_{\tilde{x}} \underline{a} \geq 0$$

for all $\underline{a}$. It will be positive definite if and only if $E(|\tilde{y}|^2) > 0$ or $\tilde{y} \neq 0$ for any $\underline{a}$. Thus, if any random variable (element of $\tilde{x}$) can be expressed as a linear combination of the others, $\underline{C}_{\tilde{x}}$ will only be positive semidefinite.

$$3) \quad \underline{C}_{\tilde{x}} = \frac{1}{2} \begin{bmatrix} \underline{A} & -\underline{B} \\ \underline{B} & \underline{A} \end{bmatrix} \quad \text{where } \underline{A} = \begin{bmatrix} 4 & 2 \\ 2 & 4 \end{bmatrix}$$

$$\underline{B} = \begin{bmatrix} 0 & 2 \\ -2 & 0 \end{bmatrix}$$

Note that $\underline{A}^T = \underline{A}$, $\underline{B}^T = -\underline{B}$ as required.

$$\Rightarrow \underline{C}_{\tilde{x}} = \underline{A} + j \underline{B} = \begin{bmatrix} 4 & 2 + 2j \\ 2 - 2j & 4 \end{bmatrix}$$

The complex Gaussian vector $\tilde{x} = \begin{bmatrix} u_1 + jv_1 \\ u_2 + jv_2 \end{bmatrix}$ will have the covariance matrix $\underline{C}_{\tilde{x}}$.

$$4) \quad a) \quad \underline{x}^{-1} = \frac{1}{2} \begin{bmatrix} 2 & -1 & 0 & 1 \\ -1 & 2 & -1 & 0 \\ 0 & -1 & 2 & -1 \\ 1 & 0 & -1 & 2 \end{bmatrix}$$

$$\underline{C} \underline{\tilde{x}}^{-1} = \frac{1}{4} \begin{bmatrix} 2 & -1+j \\ -1+j & 2 \end{bmatrix}$$

Let $\underline{x} = (\underline{u}^T \underline{v}^T)^T$ where $\underline{u} = (u_1, u_2)^T$, $\underline{v} = (v_1, v_2)^T$

$$2 \underline{x}^T \underline{C} \underline{\tilde{x}}^{-1} \underline{x} = (\underline{u}^T \underline{v}^T) \left[\begin{array}{cc|cc} 2 & -1 & 0 & 1 \\ -1 & 2 & -1 & 0 \\ \hline 0 & -1 & 2 & -1 \\ 1 & 0 & -1 & 2 \end{array} \right] \begin{bmatrix} \underline{u} \\ \underline{v} \end{bmatrix}$$

$$= (\underline{u}^T \underline{v}^T) \begin{bmatrix} \underline{C} & \underline{D}^T \\ \underline{D} & \underline{C} \end{bmatrix} \begin{bmatrix} \underline{u} \\ \underline{v} \end{bmatrix}$$

$$= \underline{u}^T \underline{C} \underline{u} + \underline{v}^T \underline{C} \underline{v} + \underline{u}^T \underline{D}^T \underline{v} + \underline{v}^T \underline{D} \underline{u}$$

$$= \underline{u}^T \underline{C} \underline{u} + \underline{v}^T \underline{C} \underline{v} + \underline{v}^T \underline{D} \underline{u} + \underline{v}^T \underline{D} \underline{u}$$

$$= \underline{u}^T \underline{C} \underline{u} + \underline{v}^T \underline{D} \underline{v} + 2 \underline{v}^T \underline{D} \underline{u}$$

$$\underline{\tilde{x}}^H \underline{C} \underline{\tilde{x}}^{-1} \underline{\tilde{x}} = (\underline{u} - j \underline{v})^T \frac{1}{4} (\underline{C} + j \underline{D}) (\underline{u} + j \underline{v})$$

$$= \frac{1}{4} (\underline{u} - j \underline{v})^T (\underline{C} \underline{u} + j \underline{C} \underline{v} + j \underline{D} \underline{u} - \underline{D} \underline{v})$$

$$= \frac{1}{4} (\underline{u}^T \underline{C} \underline{u} + j \underline{u}^T \underline{C} \underline{v} + j \underline{u}^T \underline{D} \underline{u} - \underline{u}^T \underline{D} \underline{v} - j \underline{v}^T \underline{C} \underline{u} + \underline{v}^T \underline{C} \underline{v} + \underline{v}^T \underline{D} \underline{u} + j \underline{v}^T \underline{D} \underline{v})$$

But $\underline{u}^T \underline{D} \underline{u} = (\underline{u}^T \underline{D} \underline{u})^T = \underline{u}^T \underline{D}^T \underline{u} = -\underline{u}^T \underline{D} \underline{u} = 0$

since $\underline{D}^T = -\underline{D}$ and similarly for $\underline{v}^T \underline{D} \underline{v} = 0$

Also, $\underline{v}^T \underline{C} \underline{u} = \underline{u}^T \underline{C} \underline{v}$ since $\underline{C}^T = \underline{C}$

$$\begin{aligned}
 \underline{\tilde{x}}^H \underline{C} \underline{\tilde{x}} &= 1/4 (\underline{u}^T \underline{C} \underline{u} - \underline{u}^T \underline{D} \underline{v} + \underline{v}^T \underline{C} \underline{v} + \underline{v}^T \underline{D} \underline{u}) \\
 &= 1/4 (\underline{u}^T \underline{C} \underline{u} - \underline{v}^T \underline{D}^T \underline{u} + \underline{v}^T \underline{C} \underline{v} + \underline{v}^T \underline{D} \underline{u}) \\
 &= 1/4 (\underbrace{\underline{u}^T \underline{C} \underline{u} + \underline{v}^T \underline{C} \underline{v} + 2 \underline{v}^T \underline{D} \underline{u}}_{2 \underline{\tilde{x}}^T \underline{C} \underline{\tilde{x}}}) \\
 &= \frac{1}{2} \underline{\tilde{x}}^T \underline{C} \underline{\tilde{x}}
 \end{aligned}$$

$$b) \det(\underline{C}_{\underline{x}}) = 4$$

$$\det(\underline{C}_{\underline{\tilde{x}}}) = \det \begin{pmatrix} 4 & 2+2j \\ 2-2j & 4 \end{pmatrix}$$

$$= 16 - |2+2j|^2 = 16 - 8 = 8$$

$$\frac{\det^2(\underline{C}_{\underline{\tilde{x}}})}{16} = \frac{64}{16} = 4$$

$$5) m_1 = \begin{bmatrix} a-b \\ b \ a \end{bmatrix} \quad m_2 = \begin{bmatrix} c-d \\ d \ c \end{bmatrix}$$

$$m_1 \rightarrow a+jb \quad m_2 \rightarrow c+jd$$

$$\begin{aligned}
 a) \alpha(a+jb) &= \alpha a + j\alpha b \rightarrow \begin{bmatrix} \alpha a - \alpha b \\ \alpha b \ \alpha a \end{bmatrix} \\
 &= \alpha m_1
 \end{aligned}$$

$$b) (a+jb) + (c+jd) = (a+c) + j(b+d)$$

$$\rightarrow \begin{bmatrix} a+c & -(b+d) \\ b+d & a+c \end{bmatrix} = m_1 + m_2$$

$$c) (a+jb)(c+jd) = (ac-bd) + j(ad+bc)$$

$$\rightarrow \begin{bmatrix} ac-bd & -(ad+bc) \\ ad+bc & ac-bd \end{bmatrix} = m_1 m_2$$

6) Transform 2×2 blocks or let

$$\underline{A} = \begin{bmatrix} m_1 \\ m_2 \\ m_3 \end{bmatrix} \rightarrow \begin{bmatrix} a_1 + j b_1 \\ a_2 + j b_2 \\ a_3 + j b_3 \end{bmatrix}$$

$$\text{But } \underline{A}^T \underline{A} = m_1^T m_1 + m_2^T m_2 + m_3^T m_3$$

(Note that if $m \rightarrow a + j b$, $m^T \rightarrow a - j b$)

$$\begin{aligned} \Rightarrow \underline{A}^T \underline{A} &\rightarrow (a_1 - j b_1)(a_1 + j b_1) \\ &\quad + (a_2 - j b_2)(a_2 + j b_2) \\ &\quad + (a_3 - j b_3)(a_3 + j b_3) \\ &= a_1^2 + b_1^2 + a_2^2 + b_2^2 + a_3^2 + b_3^2 + j^0 \end{aligned}$$

Transforming back we have

$$\underline{A}^T \underline{A} = \begin{bmatrix} a_1^2 + b_1^2 + a_2^2 + b_2^2 + a_3^2 + b_3^2 & 0 \\ 0 & \searrow \end{bmatrix}$$

$$7) \text{ Since } E(\tilde{x}^{*3}) = E(\tilde{x}^3)^*$$

$$E(\tilde{x}^* \tilde{x}^2) = E(\tilde{x} \tilde{x}^{*2})^*$$

we need only show that $E(\tilde{x}^3) = 0$, $E(\tilde{x}^* \tilde{x}^2) = 0$.

$$\text{But } \frac{\partial^3 \phi_{\tilde{x}}(\tilde{w})}{\partial \tilde{w}^{*3}} = (1/2)^3 E(\tilde{x}^3)$$

(see development of (15B.3) in App 15B)

$$\phi_{\tilde{w}}(\tilde{w}) = e^{-1/4 \sigma^2 |\tilde{w}|^2}$$

$$\frac{\partial \phi}{\partial \tilde{w}^*} = e^{-1/4 \sigma^2 |\tilde{w}|^2} (-1/4 \sigma^2 \tilde{w})$$

$$\frac{\partial^2 \phi}{\partial \tilde{w}^{*2}} = e^{-1/4 \sigma^2 |\tilde{w}|^2} (-1/4 \sigma^2 \tilde{w})^2$$

$$\frac{\partial^3 \phi}{\partial \tilde{w}^{*3}} = e^{-1/4 \sigma^2 |\tilde{w}|^2} (-1/4 \sigma^2 \tilde{w})^3$$

$$\Rightarrow \left. \frac{\partial^3 \phi}{\partial \tilde{w}^{*3}} \right|_{\tilde{w}=0} = 0 \Rightarrow E(\tilde{x}^3) = 0$$

Similarly,

$$E(\tilde{x}^* \tilde{x}^2) = (1/2)^3 \left. \frac{\partial^3 \phi_{\tilde{x}}(\tilde{w})}{\partial \tilde{w} \partial \tilde{w}^{*2}} \right|_{\tilde{w}=0}$$

$$\frac{\partial^2 \phi}{\partial \tilde{w}^{*2}} = e^{-1/4 \sigma^2 |\tilde{w}|^2} (-1/4 \sigma^2 \tilde{w})^2$$

$$\begin{aligned} \frac{\partial^3 \phi}{\partial \tilde{w} \partial \tilde{w}^{*2}} &= e^{-1/4 \sigma^2 |\tilde{w}|^2} \left(\frac{1}{16} \sigma^4 \tilde{w} \right) \\ &\quad + (-1/4 \sigma^2 \tilde{w})^2 e^{-1/4 \sigma^2 |\tilde{w}|^2} (-1/4 \sigma^2 \tilde{w}^*) \end{aligned}$$

which when evaluated at $\tilde{w} = 0$ is zero.

$$8) E(\tilde{x}^2) = 0$$

$$\Rightarrow E((u+jv)(u+jv)) = 0$$

$$E(u^2) - E(v^2) + j E(uv) + j E(vu) = 0$$

$$\Rightarrow E(u^2) = E(v^2), E(uv) = 0$$

$$\sigma \text{ var}(u) = \text{var}(v), \text{cov}(u, v) = 0$$

This is the usual complex Gaussian random variable assumption.

$$\begin{aligned}
 9) \quad E[(\underline{\tilde{x}} - \underline{\tilde{\mu}})(\underline{\tilde{x}} - \underline{\tilde{\mu}})^T] &= \\
 &= E\left[\left((\underline{u} - \underline{\mu}_u) + j(\underline{v} - \underline{\mu}_v)\right)\left(\quad\right)^T\right] \\
 &= E[(\underline{u} - \underline{\mu}_u)(\underline{u} - \underline{\mu}_u)^T] - E[(\underline{v} - \underline{\mu}_v)(\underline{v} - \underline{\mu}_v)^T] \\
 &\quad + j\left(E[(\underline{u} - \underline{\mu}_u)(\underline{v} - \underline{\mu}_v)^T] + E[(\underline{v} - \underline{\mu}_v)(\underline{u} - \underline{\mu}_u)^T]\right) \\
 &= \underline{0}
 \end{aligned}$$

$$\begin{aligned}
 \Rightarrow \quad \underline{C}_{uu} &= \underline{C}_{vv}, & \underline{C}_{uv} &= -\underline{C}_{vu} \\
 &= \underline{A}/2 & &= -\underline{B}/2
 \end{aligned}$$

$$\begin{aligned}
 \text{or } \underline{C}_x &= E\left[\left(\begin{bmatrix} \underline{u} \\ \underline{v} \end{bmatrix} - E\left[\begin{bmatrix} \underline{u} \\ \underline{v} \end{bmatrix}\right]\right)\left(\begin{bmatrix} \underline{u} \\ \underline{v} \end{bmatrix} - E\left[\begin{bmatrix} \underline{u} \\ \underline{v} \end{bmatrix}\right]\right)^T\right] \\
 &= \begin{bmatrix} \underline{A}/2 & -\underline{B}/2 \\ \underline{B}/2 & \underline{A}/2 \end{bmatrix}
 \end{aligned}$$

$$\begin{aligned}
 10) \quad E(\hat{\sigma}^2) &= E(\underline{\tilde{x}}^H \underline{A} \underline{\tilde{x}}) = \text{tr}(\underline{A} \underline{\sigma}^2 \underline{B}) \quad (\text{see (15.29)}) \\
 &= \sigma^2 \text{tr}(\underline{A} \underline{B}) \\
 \Rightarrow \text{tr}(\underline{A} \underline{B}) &= 1
 \end{aligned}$$

$$\begin{aligned}
 \text{var}(\hat{\sigma}^2) &= \text{tr}(\underline{A} \underline{\sigma}^2 \underline{B} \underline{A} \underline{\sigma}^2 \underline{B}) \quad (\text{see (15.30)}) \\
 &= \sigma^4 \text{tr}((\underline{A} \underline{B})^2)
 \end{aligned}$$

$$\begin{aligned}
 \text{But } \text{tr}(\underline{A} \underline{B}) &= \sum_{i=1}^N \lambda_i = 1 \\
 \text{tr}((\underline{A} \underline{B})^2) &= \sum_{i=1}^N \lambda_i^2
 \end{aligned}$$

Using the constraint we have to minimize

$$\text{tr}(\underline{A}\underline{B})^2 = \sum_{i=2}^N \lambda_i^2 + \left(1 - \sum_{i=2}^N \lambda_i\right)^2$$

$$\frac{\partial \text{tr}(\underline{A}\underline{B})^2}{\partial \lambda_k} = 2\lambda_k + 2\left(1 - \sum_{i=2}^N \lambda_i\right)(-1) = 0$$

$$k = 2, 3, \dots, N$$

$$\Rightarrow \lambda_k = 1 - \sum_{i=2}^N \lambda_i$$

Since $\sum_{i=2}^N \lambda_i$ is a constant, λ_k must be the same for all k . Thus,

$$\lambda_k = 1/N \text{ for } \text{tr}(\underline{A}\underline{B}) = 1.$$

Since

$$\begin{aligned} \underline{V}^T \underline{A} \underline{B} \underline{V} &= \underline{\Lambda} \\ &= 1/N \underline{I} \end{aligned}$$

$\underline{V}$ = modal matrix
 $\underline{\Lambda}$ = eigenvalue matrix

$$\Rightarrow \underline{A}\underline{B} = 1/N \underline{V}^T \underline{V}^{-1} = 1/N \underline{I}$$

$$\Rightarrow \underline{A} = 1/N \underline{B}^{-1} \text{ or } \hat{\sigma}^2 = 1/N \tilde{\underline{x}}^H \underline{B}^{-1} \tilde{\underline{x}}$$

$$\text{If } \underline{B} = \underline{I}, \quad \hat{\sigma}^2 = 1/N \tilde{\underline{x}}^H \tilde{\underline{x}}.$$

$$11) \quad \sum_{n=0}^{N-1} |\tilde{x}[n]|^2 = \tilde{\underline{x}}^H \tilde{\underline{x}}$$

$$E(\tilde{\underline{x}}^H \tilde{\underline{x}}) = N \sigma^2$$

$$\text{var}(\tilde{\underline{x}}^H \tilde{\underline{x}}) = \text{tr}(\underline{C}_{\tilde{\underline{x}}} \underline{C}_{\tilde{\underline{x}}})$$

(see (15.30))

$$= \text{tr}(\sigma^4 \underline{I})$$

$$= N \sigma^4$$

$$\Rightarrow \hat{\sigma}^2 = 1/N \tilde{\mathbf{x}}^H \tilde{\mathbf{x}}$$

$$\text{where } \text{var}(\hat{\sigma}^2) = \sigma^4/N$$

For real WGN, $\hat{\sigma}^2 = 1/N \mathbf{x}^T \mathbf{x}$ and $\text{var}(\hat{\sigma}^2) = 2\sigma^4/N$. Now, since $\tilde{\mathbf{x}}(n) \sim \text{CN}(0, \sigma^2)$, $\tilde{\mathbf{x}}(n) = u(n) + jv(n)$, where $u(n) \sim N(0, \sigma^2/2)$, $v(n) \sim N(0, \sigma^2/2)$. Hence, for the complex case $\hat{\sigma}^2 = 1/N \sum_{n=0}^{N-1} (u^2(n) + v^2(n))$ and we are averaging $2N$ independent samples as opposed to only N for the real case. Thus, $\text{var}(\hat{\sigma}^2)|_{\text{complex}} = 1/2 \text{var}(\hat{\sigma}^2)|_{\text{real}}$.

$$(2) \quad v(n) = \sum_{k=-\infty}^{\infty} h(k) u(n-k)$$

$$E(v(n)) = \sum_k h(k) E(u(n-k)) = 0$$

$$E(v(n)v(n+m)) = E \left[\sum_k h(k) u(n-k) \cdot \sum_l h(l) u(n+m-l) \right]$$

$$= \sum_k \sum_l h(k) h(l) \underbrace{E(u(n-k)u(n+m-l))}_{\gamma_{uu}(m+k-l)}$$

doesn't depend on $n \Rightarrow \text{WSS}$

Also, since $u(n)$ is Gaussian and $v(n)$ is the result of a linear transformation, $v(n)$ is Gaussian (u, v are jointly Gaussian)

To verify the ACF relationships

$$P_{vv}(f) = |H(f)|^2 P_{uu}(f) \\ = 1 \cdot P_{uu}(f)$$

$$\Rightarrow \Gamma_{uu}(k) = \Gamma_{vv}(k)$$

$$\text{also, } P_{uv}(f) = H(f) P_{uu}(f) \\ = -j P_{uu}(f) \quad f \geq 0 \\ j P_{uu}(f) \quad f < 0$$

$$\text{But } P_{vu}(f) = P_{uv}^*(f),$$

$$\Rightarrow P_{vu}(f) = j P_{uu}(f) \quad f \geq 0 \\ -j P_{uu}(f) \quad f < 0$$

$$= -P_{uv}(f)$$

and thus, $\Gamma_{uv}(k) = -\Gamma_{vu}(k)$. The PSD of $\tilde{x}(n)$ is from (15.3).

$$P_{\tilde{x}\tilde{x}}(f) = 2 (P_{uu}(f) + j P_{uv}(f)) \\ = 2 (P_{uu}(f) + j (-j P_{uu}(f))) \quad f \geq 0 \\ 2 (P_{uu}(f) + j (j P_{uu}(f))) \quad f < 0 \\ = 4 P_{uu}(f) \quad f \geq 0 \\ 0 \quad f < 0$$

$$(b) \quad \frac{\partial \phi^*}{\partial \theta} = \frac{1}{2} \left(\frac{\partial}{\partial \alpha} - j \frac{\partial}{\partial \beta} \right) (\alpha - j \beta)$$

$$= \frac{1}{2} \left(\frac{\partial \alpha}{\partial \alpha} - j \frac{\partial \beta}{\partial \alpha} - j \frac{\partial \alpha}{\partial \beta} - \frac{\partial \beta}{\partial \beta} \right)$$

$$= \frac{1}{2} (1 - 0 - 0 - 1) = 0$$

$$* \quad \frac{\partial}{\partial \theta} = \frac{\partial}{\partial \alpha} + j \frac{\partial}{\partial \beta}$$

$$\frac{\partial \theta}{\partial \theta} = \left(\frac{\partial}{\partial \alpha} + j \frac{\partial}{\partial \beta} \right) (\alpha + j\beta)$$

$$= 1 - 1 = 0$$

$$\frac{\partial \theta^*}{\partial \theta} = \left(\frac{\partial}{\partial \alpha} + j \frac{\partial}{\partial \beta} \right) (\alpha - j\beta)$$

$$= 1 + 1 = 2$$

$$14) \quad \frac{\partial \theta^H \underline{b}}{\partial \theta_k} = \frac{\partial}{\partial \theta_k} \sum_l b_l \theta_l^*$$

$$= \sum_l b_l \frac{\partial \theta_l^*}{\partial \theta_k} = 0 \quad \text{all } k$$

$$\Rightarrow \frac{\partial \theta^H \underline{b}}{\partial \theta} = \underline{0}$$

15) Same proof as in deriving (15.46)

$$\frac{\partial \theta^H \underline{A} \underline{\theta}}{\partial \underline{\theta}} = \underline{A}^T \underline{\theta}^*$$

Since $\underline{A}^H \neq \underline{A}$, $\frac{\partial \theta^H \underline{A} \underline{\theta}}{\partial \underline{\theta}} = (\underline{A}^H \underline{\theta})^*$

16) To find LSE

$$J = \sum_{n=0}^{N-1} |\tilde{x}[n] - \tilde{A} y^n|^2$$

From Example 15.2

$$\begin{aligned}\hat{\tilde{A}} &= \frac{\sum_{n=0}^{N-1} \tilde{x}(n) \tilde{s}^*[n]}{\sum_{n=0}^{N-1} |\tilde{s}(n)|^2} \\ &= \frac{\sum_{n=0}^{N-1} \tilde{x}(n) r^{*n}}{\sum_{n=0}^{N-1} |r|^{2n}}\end{aligned}$$

MLE will be the same.

$$\begin{aligned}E(\hat{\tilde{A}}) &= \frac{\sum_n E(\tilde{x}(n)) r^{*n}}{\sum_n |r|^{2n}} \\ &= \frac{\tilde{A} \sum_n r^n r^{*n}}{\sum_n |r|^{2n}} = \tilde{A}\end{aligned}$$

$$\begin{aligned}\text{var}(\hat{\tilde{A}}) &= \frac{\sum_n \text{var}(\tilde{x}(n)) |r|^{2n}}{(\sum_n |r|^{2n})^2} \\ &= \frac{\sigma^2}{\sum_n |r|^{2n}}\end{aligned}$$

$$\begin{aligned}\text{If } |r| < 1, \quad \text{var}(\hat{\tilde{A}}) &\rightarrow \frac{\sigma^2}{(1-|r|^2)^{-1}} \\ &= \sigma^2 (1-|r|^2)\end{aligned}$$

$$\text{If } |r| \geq 1, \quad \text{var}(\hat{\tilde{A}}) \rightarrow 0$$

17) Equivalently, we have

$$\tilde{x}(n) = A + \tilde{w}(n) \quad \text{or}$$

$$u(n) = A + w_R(n)$$

$$v(n) = w_I(n)$$

Now let $\underline{x} = \begin{bmatrix} u \\ v \end{bmatrix}$ and $\underline{w} = \begin{bmatrix} w_R \\ w_I \end{bmatrix}$ so that

$$\underline{x} = A \begin{bmatrix} 1 \\ 0 \end{bmatrix} + \underline{w}$$

where $E(\underline{w}) = \underline{0}$, $\underline{C}_w = \sigma^2/2 \underline{I} = \underline{C}_x$

The real BLUE can now be used.

$$\begin{aligned} \hat{A} &= \frac{[1^T \ 0^T] \underline{C}_w^{-1} \underline{x}}{[1^T \ 0^T] \underline{C}_w^{-1} \begin{bmatrix} 1 \\ 0 \end{bmatrix}} \\ &= \frac{1}{N} \sum_{n=0}^{N-1} u(n) = \frac{1}{N} \sum_{n=0}^{N-1} \text{Re}(\tilde{x}(n)) \end{aligned}$$

In Example 15.7 $\tilde{A}$ is complex and if $\underline{C} = \sigma^2 \underline{I}$, which conforms to the assumptions for this problem,

$$\hat{\tilde{A}} = \frac{1^T \tilde{\underline{x}}}{1^T 1} = \frac{1}{N} \sum_{n=0}^{N-1} \tilde{x}(n)$$

18) From (15.52)

$$\begin{aligned} \mathcal{I}(\theta) &= 2 \text{Re} \left[\frac{\partial \tilde{u}^H(\theta)}{\partial \theta} \underline{\tilde{x}}^{-1} \frac{\partial \tilde{u}(\theta)}{\partial \theta} \right] \\ &= \frac{2}{\sigma^2} \text{Re} \left[\sum_{n=0}^{N-1} \left| \frac{\partial (\tilde{u}(\theta))_n}{\partial \theta} \right|^2 \right] \end{aligned}$$

$$= \frac{2}{\sigma^2} \sum_{n=0}^{N-1} \left| \frac{\partial \tilde{S}(n; \theta)}{\partial \theta} \right|^2$$

$$\text{var}(\hat{\theta}) \geq \frac{\sigma^2/2}{\sum_{n=0}^{N-1} \left(\frac{\partial \text{Re}[\tilde{S}(n; \theta)]}{\partial \theta} \right)^2 + \left(\frac{\partial \text{Im}[\tilde{S}(n; \theta)]}{\partial \theta} \right)^2}$$

For real data (see (3.14))

$$\text{var}(\hat{\theta}) \geq \frac{\sigma^2}{\sum_{n=0}^{N-1} \left(\frac{\partial S(n; \theta)}{\partial \theta} \right)^2}$$

In complex case there is information in both the real and imaginary parts of the signal. Also, the $\sigma^2/2$ factor accounts for the variance of each real noise sample.

19) From (15.52) with $\underline{\tilde{x}} = \sigma^2 \underline{I}$

$$\begin{aligned} I(\sigma^2) &= \text{tr} \left[\underline{\tilde{x}}^{-1} \frac{\partial \underline{\tilde{x}}}{\partial \sigma^2} \underline{\tilde{x}}^{-1} \frac{\partial \underline{\tilde{x}}}{\partial \sigma^2} \right] \\ &= \text{tr} \left[\left(\frac{1}{\sigma^2} \underline{I} \right) (\underline{I}) \left(\frac{1}{\sigma^2} \underline{I} \right) (\underline{I}) \right] \\ &= N / \sigma^4 \end{aligned}$$

$$\text{var}(\hat{\sigma}^2) \geq \sigma^4 / N$$

$$p(\underline{\tilde{x}}; \sigma^2) = \frac{1}{\pi^N \det(\sigma^2 \underline{I})} e^{-\frac{1}{\sigma^2} \underline{\tilde{x}}^H \underline{\tilde{x}}}$$

$$\begin{aligned}
\frac{\partial \ln p}{\partial \sigma^2} &= - \frac{\partial \ln \det(\sigma^2 \mathbf{I})}{\partial \sigma^2} - \frac{\partial}{\partial \sigma^2} \left(\frac{1}{\sigma^2} \tilde{\mathbf{X}}^H \tilde{\mathbf{X}} \right) \\
&= - \frac{\partial \ln \sigma^{2N}}{\partial \sigma^2} - \frac{\partial}{\partial \sigma^2} \left(\frac{1}{\sigma^2} \tilde{\mathbf{X}}^H \tilde{\mathbf{X}} \right) \\
&= -N/\sigma^2 + 1/\sigma^4 \tilde{\mathbf{X}}^H \tilde{\mathbf{X}} \\
&= N/\sigma^4 \left(\frac{1}{N} \sum_{n=0}^{N-1} |\tilde{x}(n)|^2 - \sigma^2 \right)
\end{aligned}$$

efficient estimator

$$\begin{aligned}
20) \quad \hat{\theta}_i &= \underline{a}_i^H \tilde{\mathbf{x}} \Rightarrow \hat{\underline{\theta}} = \underline{A} \tilde{\mathbf{x}} \quad \underline{A} = \begin{bmatrix} \underline{a}_1^H \\ \vdots \\ \underline{a}_p^H \end{bmatrix} \\
E(\hat{\underline{\theta}}) &= \underline{\theta} \Rightarrow E(\underline{A}(\underline{H}\underline{\theta} + \tilde{\mathbf{w}})) \\
&= \underline{A}\underline{H}E(\underline{\theta}) = \underline{\theta}
\end{aligned}$$

$$\text{or } \underline{A}\underline{H} = \mathbf{I}$$

$$\begin{aligned}
\text{var}(\hat{\theta}_i) &= E(|\theta_i - \hat{\theta}_i|^2) \\
&= E(|\underline{a}_i^H \tilde{\mathbf{x}} - \underline{a}_i^H E(\tilde{\mathbf{x}})|^2) \\
&= E[\underline{a}_i^H (\tilde{\mathbf{x}} - E(\tilde{\mathbf{x}})) (\tilde{\mathbf{x}} - E(\tilde{\mathbf{x}}))^H \underline{a}_i] \\
&= \underline{a}_i^H \underline{\zeta} \underline{a}_i
\end{aligned}$$

$$\underline{A}\underline{H} = \mathbf{I} \Rightarrow \begin{bmatrix} \underline{a}_1^H \\ \vdots \\ \underline{a}_p^H \end{bmatrix} \underline{H} = \mathbf{I}$$

$$\text{or } \underline{H}^H [\underline{a}_1 \dots \underline{a}_p] = \mathbf{I}$$

$$\underline{H}^H \underline{a}_i = \underline{e}_i \quad \text{where } \underline{e}_i = [0 \dots 0 \underset{\uparrow i\text{th place}}{1} 0 \dots 0]^T$$

Using (15.51) we have

$$\underline{a}_{i\text{opt}} = \underline{\zeta}^{-1} \underline{H} (\underline{H}^H \underline{\zeta}^{-1} \underline{H})^{-1} \underline{e}_i$$

$$\hat{\theta}_i = \underline{a}_{i \text{ opt}} \tilde{\underline{x}} = \underline{e}_i^H (\underline{H} \underline{C}^{-1} \underline{H})^{-1} \underline{H}^H \underline{C}^{-1} \tilde{\underline{x}}$$

$$\Rightarrow \hat{\underline{\theta}} = (\underline{H} \underline{C} \underline{H})^{-1} \underline{H}^H \underline{C}^{-1} \tilde{\underline{x}}$$

$$21) \quad \tilde{\underline{x}} = \tilde{\underline{A}}_1 \underline{e}_1 + \tilde{\underline{A}}_2 \underline{e}_2 + \tilde{\underline{w}} = \underline{E} \tilde{\underline{A}} + \tilde{\underline{w}}$$

$$\text{where } \underline{E} = (\underline{e}_1, \underline{e}_2), \quad \tilde{\underline{A}} = (\tilde{\underline{A}}_1, \tilde{\underline{A}}_2)^T$$

The MLE will minimize

$$\begin{aligned} J(\tilde{\underline{A}}, \underline{f}) &= (\tilde{\underline{x}} - \underline{E} \tilde{\underline{A}})^H (\tilde{\underline{x}} - \underline{E} \tilde{\underline{A}}) \\ &= \tilde{\underline{x}}^H \tilde{\underline{x}} - \tilde{\underline{x}}^H \underline{E} \tilde{\underline{A}} - \tilde{\underline{A}}^H \underline{E}^H \tilde{\underline{x}} + \tilde{\underline{A}}^H \underline{E}^H \underline{E} \tilde{\underline{A}} \end{aligned}$$

Taking the complex gradient we have

$$\frac{\partial J}{\partial \tilde{\underline{A}}} = \underline{0} - \frac{\partial}{\partial \tilde{\underline{A}}} \tilde{\underline{x}}^H \underline{E} \tilde{\underline{A}} - \underline{0} + \frac{\partial}{\partial \tilde{\underline{A}}} \tilde{\underline{A}}^H \underline{E}^H \underline{E} \tilde{\underline{A}}$$

Using (15.44) and (15.46)

$$\frac{\partial J}{\partial \tilde{\underline{A}}} = -(\underline{E}^H \tilde{\underline{x}})^* + (\underline{E}^H \underline{E} \tilde{\underline{A}})^* = \underline{0}$$

$$\Rightarrow \hat{\underline{A}} = (\underline{E}^H \underline{E})^{-1} \underline{E}^H \tilde{\underline{x}}$$

$$\begin{aligned} J(\hat{\underline{A}}, \underline{f}) &= (\tilde{\underline{x}} - \underline{E}(\underline{E}^H \underline{E})^{-1} \underline{E}^H \tilde{\underline{x}})^H \\ &\quad \cdot (\tilde{\underline{x}} - \underline{E}(\underline{E}^H \underline{E})^{-1} \underline{E}^H \tilde{\underline{x}}) \end{aligned}$$

$$= \tilde{\underline{x}}^H (\underline{I} - \underline{E}(\underline{E}^H \underline{E})^{-1} \underline{E}^H) \tilde{\underline{x}}$$

Since $\underline{I} - \underline{E}(\underline{E}^H \underline{E})^{-1} \underline{E}^H$ is idempotent

To minimize over $\underline{f}$ we need to maximize

$$J(\underline{f}) = \tilde{\underline{X}}^H \underline{E} (\underline{E}^H \underline{E})^{-1} \underline{E}^H \tilde{\underline{X}}$$

This will require a 2-D search. An approximate MLE for $|f_1 - f_2| \gg 1/N$ is found as follows:

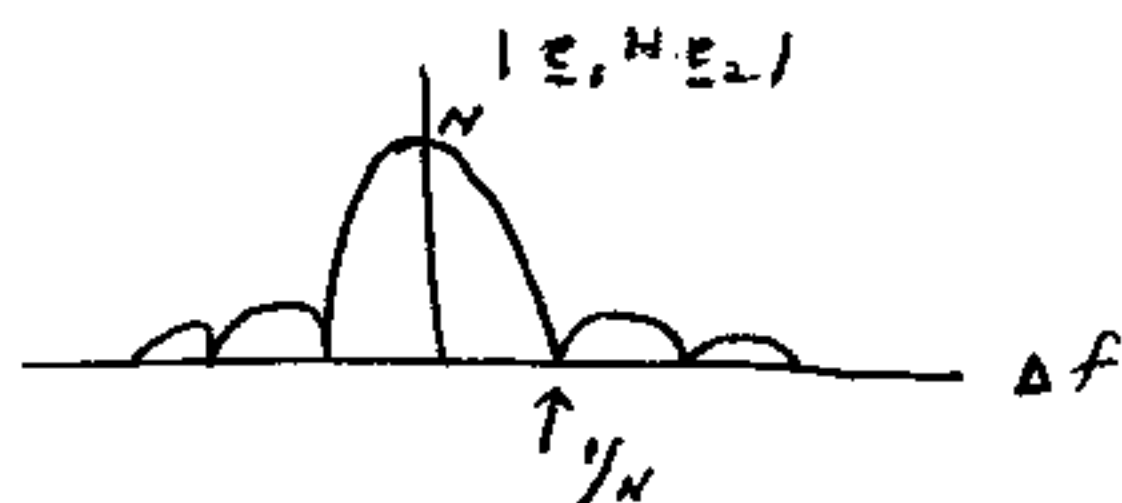
$$\underline{E}^H \underline{E} = \begin{bmatrix} \underline{e}_1^H \\ \underline{e}_2^H \end{bmatrix} \begin{bmatrix} \underline{e}_1 & \underline{e}_2 \end{bmatrix} = \begin{bmatrix} N & \underline{e}_1^H \underline{e}_2 \\ \underline{e}_2^H \underline{e}_1 & N \end{bmatrix}$$

$$\begin{aligned} \underline{e}_1^H \underline{e}_2 &= \sum_{n=0}^{N-1} e^{-j2\pi f_1 n} e^{j2\pi f_2 n} \\ &= \sum_{n=0}^{N-1} e^{j2\pi \Delta f n} \quad \Delta f = f_2 - f_1 \end{aligned}$$

$$= \frac{1 - e^{j2\pi \Delta f N}}{1 - e^{j2\pi \Delta f}}$$

$$= \frac{e^{j\pi \Delta f N}}{e^{j\pi \Delta f}} \frac{e^{-j\pi \Delta f N} - e^{j\pi \Delta f N}}{e^{-j\pi \Delta f} - e^{j\pi \Delta f}}$$

$$= e^{j\pi \Delta f (N-1)} \frac{\sin \pi \Delta f N}{\sin \pi \Delta f}$$



For $|\Delta f| \gg 1/N$, $\underline{e}_1^H \underline{e}_2 \ll N$ and thus

$$\underline{E}^H \underline{E} \approx N \underline{I}$$

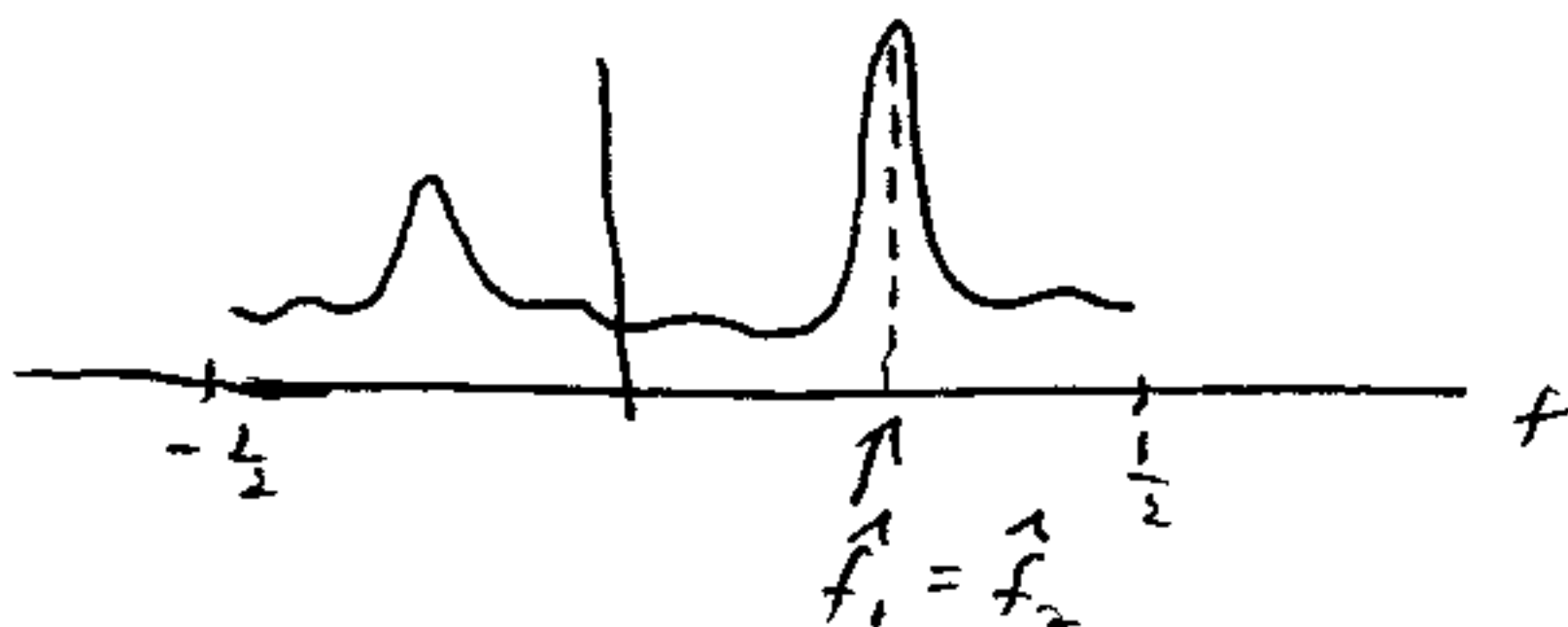
$$J(\underline{f}) \approx \frac{1}{N} \tilde{\underline{X}}^H \underline{E} \underline{E}^H \tilde{\underline{X}} = \frac{1}{N} \|\underline{E}^H \tilde{\underline{X}}\|^2$$

$$= \frac{1}{N} (|\underline{e}_1^H \tilde{\mathbf{x}}|^2 + |\underline{e}_2^H \tilde{\mathbf{x}}|^2)$$

$$= \frac{1}{N} \left| \sum_{n=0}^{N-1} \tilde{x}(n) e^{-j2\pi f_1 n} \right|^2 \\ + \frac{1}{N} \left| \sum_{n=0}^{N-1} \tilde{x}(n) e^{-j2\pi f_2 n} \right|^2$$

Approximate MLE finds peak locations of the Two largest peaks of a single periodogram, with the constraint $|f_1 - f_2| \gg 1/N$.

Without constraint we could have



22) An MLE will minimize

$$\sum_{n=0}^{N-1} |\tilde{x}(n) - \tilde{A} \tilde{s}(n)|^2$$

where $\tilde{s}(n) = e^{j2\pi(f_0 n + 1/2 n^2)}$

This is a linear LS problem with respect to $\tilde{A}$. Thus, from Example 15.2

$$\hat{\tilde{A}} = \frac{\sum_n \tilde{x}(n) \tilde{s}^*(n)}{\sum_n |\tilde{s}(n)|^2}$$

Substituting for $\tilde{A}$

$$\begin{aligned}
 & \sum_{n=0}^{N-1} (\tilde{x}[n] - \hat{\tilde{A}} \tilde{s}[n])^* (\tilde{x}[n] - \hat{\tilde{A}} \tilde{s}[n]) \\
 &= \sum_{n=0}^{N-1} \tilde{x}^*[n] (\tilde{x}[n] - \hat{\tilde{A}} \tilde{s}[n]) \\
 &\quad - \underbrace{\hat{\tilde{A}} \sum_{n=0}^{N-1} \tilde{s}^*[n] (\tilde{x}[n] - \hat{\tilde{A}} \tilde{s}[n])}_{=0} \\
 &= \sum_n |\tilde{x}[n]|^2 - \hat{\tilde{A}} \sum_n \tilde{x}^*[n] \tilde{s}[n] \\
 &= \sum_n |\tilde{x}[n]|^2 - \frac{\left| \sum_n \tilde{x}[n] \tilde{s}^*[n] \right|^2}{\sum_n |\tilde{s}[n]|^2} \\
 &= \sum_n |\tilde{x}[n]|^2 - \frac{\left| \sum_n \tilde{x}[n] e^{-j2\pi(f_0 n + \frac{1}{2}\alpha n^2)} \right|^2}{N}
 \end{aligned}$$

Hence, we need to maximize over f_0, α

$$\left| \sum_{n=0}^{N-1} \tilde{x}[n] e^{-j2\pi(f_0 n + \frac{1}{2}\alpha n^2)} \right|^2. \quad \text{To}$$

compute this efficiently let $\tilde{y}[n] = \tilde{x}[n] e^{-j\pi\alpha n^2}$

$$\Rightarrow \left| \sum_{n=0}^{N-1} \tilde{y}[n] e^{-j2\pi f_0 n} \right|^2, \quad \text{can use}$$

FFT on $\tilde{y}[n]$ for each α .

$$23) \quad \hat{\alpha} = E(\alpha | \underline{u}, \underline{v}) = \int \alpha p(\alpha | \underline{u}, \underline{v}) d\alpha$$

$$\hat{\beta} = E(\beta | \underline{u}, \underline{v}) = \int \beta p(\beta | \underline{u}, \underline{v}) d\beta$$

$$\text{or } \hat{\alpha} = \iint \alpha p(\alpha, \beta | \underline{u}, \underline{v}) d\alpha d\beta$$

$$\hat{\beta} = \iint \beta p(\alpha, \beta | \underline{u}, \underline{v}) d\alpha d\beta$$

$$\hat{\theta} = \hat{\alpha} + g\hat{\beta} = \iint (\alpha + g\beta) p(\alpha, \beta | \underline{u}, \underline{v}) d\alpha d\beta$$

$$= \iint \theta p(\alpha, \beta | \underline{u}, \underline{v}) d\alpha d\beta$$

$$= E(\theta | \underline{u}, \underline{v}) = E(\underline{\theta} | \underline{x})$$

Note that $p(\underline{\theta} | \underline{x})$ is just $p(\alpha, \beta | \underline{u}, \underline{v})$ in disguise.

24) From (15.52) with $\underline{\tilde{x}} = P_0 \underline{Q}$ where $\underline{Q}$ is the covariance matrix corresponding to a process with PSD $Q(f)$,

$$\mathcal{I}(P_0) = \text{tr} \left\{ \underline{\tilde{x}}^{-1} \frac{\partial \underline{\tilde{x}}}{\partial P_0} \underline{\tilde{x}}^{-1} \frac{\partial \underline{\tilde{x}}}{\partial P_0} \right\}$$

$$= \text{tr} \left\{ \frac{1}{P_0} \underline{Q}^{-1} \underline{Q} \frac{1}{P_0} \underline{Q}^{-1} \underline{Q} \right\}$$

$$= \frac{1}{P_0^2} \text{tr}(\underline{I}) = N/P_0^2$$

$$\Rightarrow \text{var}(\hat{P}_0) \geq P_0^2/N$$

From (15.68)

$$\mathcal{I}(P_0) = N \int_{-\frac{1}{2}}^{\frac{1}{2}} \left(\frac{\partial \ln P_{\tilde{x}}(f)}{\partial P_0} \right)^2 df$$

$$\begin{aligned}
 &= N \int_{-\frac{1}{2}}^{\frac{1}{2}} \left(\frac{\partial \ln p_0 \phi(t)}{\partial p_0} \right)^2 dt \\
 &= N \int_{-\frac{1}{2}}^{\frac{1}{2}} \frac{1}{p_0^2} dt = N/p_0^2
 \end{aligned}$$

Same result (not true in general)

$$\begin{aligned}
 25) \quad x(f_k) &= \sum_n \tilde{x}(n) e^{-j2\pi f_k n} \quad (f_k = k/N) \\
 &= \sum_n (\tilde{A} e^{j2\pi f_c n} + \tilde{w}(n)) e^{-j2\pi f_k n} \\
 &= \sum_n \tilde{w}(n) e^{-j2\pi f_k n} \quad f_k \neq f_c \\
 &\quad \sum_n (\tilde{A} + \tilde{w}(n) e^{-j2\pi f_k n}) \quad f_k = f_c \\
 &= W(f_k) \quad k \neq l \\
 &\quad N\tilde{A} + W(f_k) \quad k = l
 \end{aligned}$$

Where $W(f_k)$ is the DFT of $\tilde{w}(n)$.

But $\underline{w} \sim \mathcal{CN}(\underline{0}, N\sigma^2 \underline{I})$ from Sect 15.9

$$\Rightarrow \underline{x} \sim \mathcal{CN}\left(\begin{bmatrix} 0 \\ \vdots \\ 0 \\ N\tilde{A} \\ 0 \\ \vdots \\ 0 \end{bmatrix}, N\sigma^2 \underline{I}\right)$$

Where $\underline{x} = [x(f_0) x(f_1) \dots x(f_{N-1})]^T$

$$\begin{aligned}
 \text{Now, } \text{SNR}(\text{input}) &= |E(\tilde{x}(n))|^2 / \sigma^2 \\
 &= |\tilde{A}|^2 / \sigma^2
 \end{aligned}$$

$$\text{SNR}(\text{output}) = \frac{|E(x(f_c))|^2}{\text{var}(x(f_c))}$$

$$= \frac{|N \tilde{A}|^2}{N \sigma^2} = \frac{N |\hat{A}|^2}{\sigma^2}$$

Processing gain = $10 \log_{10} N$ dB

To detect a sinusoid of unknown frequency we could form

$$\frac{1}{N} |x(f_k)|^2 = \frac{1}{N} \left| \sum_{n=0}^{N-1} \tilde{x}[n] e^{-j2\pi f_k n} \right|^2$$

and choose the maximum over f_k for comparison to a threshold. If no signal is present,

$$\begin{aligned} E \left[\frac{1}{N} |x(f_k)|^2 \right] &= \frac{1}{N} \text{var}(x(f_k)) \\ &= \sigma^2 \quad \text{for all } k \end{aligned}$$

and for a signal present

$$E \left[\frac{1}{N} |x(f_k)|^2 \right] = \frac{1}{N} \left[\text{var}(x(f_k)) + |E(x(f_k))|^2 \right]$$

$$= \frac{1}{N} (N \sigma^2 + N^2 |\tilde{A}|^2)$$

$$= \sigma^2 + N |\tilde{A}|^2 \quad \text{for } k=l$$

$$= \sigma^2 \quad \text{for } k \neq l$$

This statistic is just a sampled in frequency periodogram!