

2.6 Maximum Likelihood Estimation

- Recall two general estimation approaches
 - $\underline{\theta}$ is random, pdf $p_{\theta}(\underline{\theta})$ is known. This produces Bayes parameter estimation, special case linear mmse estimation
 - $\underline{\theta}$ is non random but unknown, or random but no information on $p_{\theta}(\underline{\theta})$. ML is a powerful method here.

- Recall MAP estimation:

choose $\hat{\underline{\theta}}$ to maximize

$$p(\hat{\underline{\theta}}|x) = \frac{p(x|\hat{\underline{\theta}})p(\hat{\underline{\theta}})}{p(x)}, \text{ equivalently } \hat{\underline{\theta}} = \arg \max_{\underline{\theta}} p(x|\underline{\theta})p(\underline{\theta})$$

2.6.1 The ML principle says — if you have no idea how $\underline{\theta}$ is distributed (if it has a distribution), then

choose $\hat{\underline{\theta}}$ to maximize $p(x|\hat{\underline{\theta}})$

Equivalent to making $p_{\theta}(\underline{\theta})$ uniform — no information.

- ML estimators may or may not be biased.
- Example: Photon Counting. Estimate arrival rate λ from count k . If we have no idea how likely different values of λ are, we use $p_{k|\lambda}(k|\lambda) = \frac{e^{-\lambda T} (\lambda T)^k}{k!}$

$$\frac{\partial}{\partial \lambda} p_{k|\lambda}(k|\lambda) = \frac{T}{k!} e^{-\lambda T} (\lambda T)^{k-1} (-\lambda T + k) \Rightarrow \hat{\lambda} = \frac{k}{T} \quad \text{ML}$$

Compare with mmse when λ exponentially distributed

$$\hat{\lambda} = \frac{b}{bT+1} (k+1) \quad \text{MMSE}$$

- Call $l(\underline{\theta}) = P_{\underline{x}|\underline{\theta}}(\underline{x}|\underline{\theta})$ the likelihood function

Since logarithm is monotonic, it is equivalent to maximize

$$L(\underline{\theta}) = \ln(l(\underline{\theta})) \quad \text{the log-likelihood (sometimes } \Lambda(\underline{\theta}))$$

The log likelihood is convenient when $p(\underline{x}|\underline{\theta})$ has an exponential form (e.g. Gaussian) or if components of $\underline{x}$ are conditionally independent

$$L(\underline{\theta}) = \ln \left[\prod_{i=1}^M P_{x_i|\underline{\theta}}(x_i|\underline{\theta}) \right] = \sum_{i=1}^M L_i(\underline{\theta})$$

- Since maxima are important, the vector of partial derivatives has a name

$$V(\underline{\theta}) = \frac{d}{d\underline{\theta}} L(\underline{\theta}) = \begin{bmatrix} \frac{\partial L(\underline{\theta})}{\partial \theta_1} \\ \vdots \\ \frac{\partial L(\underline{\theta})}{\partial \theta_K} \end{bmatrix} \quad \text{the "score"}$$

- Note that $l(\underline{\theta})$, $L(\underline{\theta})$, $V(\underline{\theta})$ all depend on $\underline{x}$, so all are random variables.

- ML estimators have some nice properties. If we make a sequence of observations $\underline{x}_i$, $i=1, \dots, N$ that are conditionally independent and identically distributed, we have asymptotic properties as $N \rightarrow \infty$:

- ML estimators are asymptotically unbiased
- ML estimators are asymptotically consistent
- ML estimators are asymptotically efficient (unbiased and meet Cramér Rao lower bound)
- ML estimators are asymptotically Gaussian

2.6.3 A very important special case leads to least squares.

Suppose we observe

$$\underline{x} = \underbrace{A \underline{\theta}}_{\text{signal}} + \underline{n}$$

where A is a known matrix $N \times M$

$\underline{\theta}$ is to be estimated M

$\underline{n}$ is white Gaussian noise. N
 $E[\underline{n} \underline{n}^T] = \sigma_n^2 \mathbf{I}$

$$l(\underline{\theta}) = p_{\underline{x}|\underline{\theta}}(\underline{x}|\underline{\theta}) = \frac{1}{(2\pi)^N \left(\frac{\sigma_n^2}{2}\right)^N} \exp \left[-\frac{1}{\sigma_n^2} (\underline{x} - A \underline{\theta})^T (\underline{x} - A \underline{\theta}) \right]$$

$$L(\underline{\theta}) \sim -\frac{1}{\sigma_n^2} |\underline{x} - A \underline{\theta}|^2 = -\frac{1}{\sigma_n^2} \sum_{i=1}^N |x_i - \underline{a}_i^T \underline{\theta}|^2 \quad \begin{array}{l} \text{minimize} \\ \text{sum of squared} \\ \text{errors.} \end{array}$$

This is a linearly parameterized signal model, and we want the best fit to observations.

Solution: $\frac{d}{d\underline{\theta}} (\underline{x} - A \underline{\theta})^T (\underline{x} - A \underline{\theta}) = -(\underline{x}^T - \underline{\theta}^T A^T) A = \underline{0}$

or $A^T A \hat{\underline{\theta}} = A^T \underline{x}$, which yields

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

"deterministic normal equations"

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

(projection of $\underline{x}$ onto subspace spanned by A cols)

$$\hat{\underline{r}} = \underline{x} - \hat{\underline{s}} = (\mathbf{I} - A (A^T A)^{-1} A^T) \underline{x}$$

(orthog complement — proj'n of $\underline{x}$ onto subspace orthog to A columns)

Deterministic normal equations:

— basis of signal space is columns of $A = [\underline{a}_1 | \underline{a}_2 | \dots | \underline{a}_N]$

— Gram matrix of inner prods $[G]_{i,k} = \underline{a}_k^T \underline{a}_i$

so $G = A^T A$

— Inner prod with observations $[p]_k = \underline{a}_k^T \underline{x}$

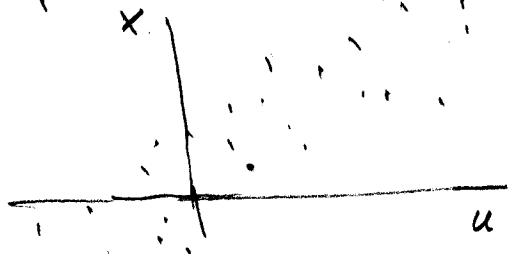
so $\underline{p} = A^T \underline{x}$

— Therefore $A^T A \hat{\underline{\theta}} = A^T \underline{x}$ is just $G \hat{\underline{\theta}} = \underline{p}$

$\hat{\underline{\theta}} = (A^T A)^{-1} A^T \underline{x}$ is just $\hat{\underline{\theta}} = G^{-1} \underline{p}$

example: best straight line

2.6.4



x vs. u

model $x = \theta_1 + \theta_2 u$

$$\underline{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_N \end{bmatrix}$$

$$A = \begin{bmatrix} 1 & u_1 \\ \vdots & \vdots \\ 1 & u_N \end{bmatrix}$$

$$\underline{\theta} = \begin{bmatrix} \theta_1 \\ \theta_2 \end{bmatrix}$$

$$\underline{x} = A \underline{\theta} + \underline{n}$$

example: channel sounding (like PSAM)

|P|P|P|P|

$$r(k) = g(k)c(k) + n(k)$$

assume g is constant but unknown (including σ_g^2) during pilots

$$r(k) = g + n(k) \quad \text{what is ML est of } g?$$

Since $n(k)$ is white, Gaussian, ML becomes least squares

$$\underline{x} = \begin{bmatrix} r(1) \\ \vdots \\ r(N) \end{bmatrix}$$

$$A = \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix}$$

$$\underline{\theta} = g$$

$$\underline{x} = A \underline{\theta} + \underline{n}$$

$$\begin{bmatrix} c(1) \\ \vdots \\ c(N) \end{bmatrix}$$

More generally, get impulse response

S N. Grier, D. Fabre & S. A. Mahmoud "Least Sum of Squared Errors Channel Estimation" *Proc IEE Part F*, vol 138 no 4 August 1991

R. A. Ziegler & J. M. Cioffi "Estimation of Time Varying Digital Radio Channels" *IEEE Trans Veh Technol* vol 41 no 2 pp 134-151, May 1992