

Outline

0. Review & course evaluation

1. Maximum likelihood estimation

2. Cramer Rao inequality

3. Bayes' estimation

0. Review

What we have done so far?

Find the best parameter estimation based on measurements &

a priori information with different principles

Supposed mathematical model: $\vec{y} = H\vec{x}$; $\vec{\tilde{y}} = H\vec{x} + \vec{v}$; $\hat{\vec{y}} = H\hat{\vec{x}}$

a) Least squares approximation

$$J = \min \frac{1}{2} (\hat{\vec{y}} - \vec{\tilde{y}})^T (\hat{\vec{y}} - \vec{\tilde{y}})$$

$$\hat{\vec{x}} = (H^T H)^{-1} H^T \vec{\tilde{y}}$$

weighted LSA

$$\hat{\vec{x}} = (H^T W H)^{-1} H^T W \vec{\tilde{y}} \quad \text{No mathematical criteria for } W$$

b) Minimum variance estimation (Probability)

$$\hat{\vec{x}} = M \vec{\tilde{y}} + \vec{n} \Rightarrow \vec{n} = 0 \quad MH = I$$

$$J = \min \frac{1}{2} \text{Tr} [E \{ (\hat{\vec{x}} - \vec{x}) (\hat{\vec{x}} - \vec{x})^T \}]$$

$$\hat{\vec{x}} = (H^T R^{-1} H)^{-1} H^T R^{-1} \vec{\tilde{y}}$$

is identical to LSA, provided that $W = R^{-1}$

Minimum variance estimation with a priori

$$\hat{\vec{x}} = M \vec{\tilde{y}} + N \hat{\vec{x}}_a + \vec{n} \Rightarrow \vec{n} = 0 \quad MH + N = I$$

$$\text{Solution } \hat{\vec{x}} = (H^T R^{-1} H + Q^{-1})^{-1} (H^T R^{-1} \vec{\tilde{y}} + Q^{-1} \hat{\vec{x}}_a)$$

1. Maximum likelihood estimation Not always biased

Recall: a Gaussian distribution $Y \sim N(\mu, \sigma^2)$

the probability density function

$$f_Y(y) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(y-\mu)^2}{2\sigma^2}\right)$$

$$E[\tilde{Y}_i] = \mu$$

Question: Given $\{\tilde{Y}_1, \tilde{Y}_2, \dots, \tilde{Y}_m\}$

$$\Rightarrow \text{Var}[\tilde{Y}_i] = \sigma^2$$

how to estimate μ and σ^2 from observed data?

$$\hat{\mu} = \frac{1}{m} \sum_{i=1}^m \tilde{Y}_i; \quad \text{Let } \vec{x} = [\mu, \sigma^2]$$

Basic idea:

A good estimate of $\vec{x}$ would be the value of $\hat{\vec{x}}$ that maximize the probability (likelihood) of getting the data we observed.

Define likelihood function

$$L(\vec{\tilde{Y}}; \vec{x}) = f(\tilde{Y}_1; \vec{x}) f(\tilde{Y}_2; \vec{x}) \cdots f(\tilde{Y}_m; \vec{x})$$
$$= \prod_{i=1}^m f(\tilde{Y}_i; \vec{x})$$

In textbook $f(\tilde{Y}_i; \vec{x})$ is written as $p(\tilde{Y}_i | \vec{x})$

For the Gaussian example:

$$L(\vec{\tilde{Y}}; \vec{x}) = \prod_{i=1}^m f(\tilde{Y}_i; \vec{x}) = \left(\frac{1}{2\pi\sigma^2}\right)^{\frac{m}{2}} \exp\left(-\frac{\sum_{i=1}^m (\tilde{Y}_i - \mu)^2}{2\sigma^2}\right)$$

The Gaussian distribution is a monotonic exponential function for μ and σ^2 .

$$\ln[L(\vec{\tilde{Y}}; \vec{x})] = -\frac{m}{2} \ln(2\pi\sigma^2) - \frac{1}{2\sigma^2} \sum_{i=1}^m (\tilde{Y}_i - \mu)^2$$

$\text{Max } L(\vec{\tilde{Y}}; \vec{x})$ is equivalent to $\text{Max } \ln[L(\vec{\tilde{Y}}; \vec{x})]$

$$J = \max \ln[L(\vec{\tilde{Y}}; \hat{\vec{x}})]$$

Necessary condition: $\frac{\partial}{\partial x} \ln[L(\vec{\tilde{Y}}; \hat{\vec{x}})] = 0$

Sufficient condition: $\frac{\partial^2}{\partial x^2} \ln[L(\vec{\tilde{Y}}; \hat{\vec{x}})] < 0$

$$\frac{\partial}{\partial \hat{\mu}} \ln[L(\vec{\tilde{Y}}; \hat{\vec{x}})] = \frac{1}{\sigma^2} \sum_{i=1}^m (\tilde{Y}_i - \hat{\mu}) = 0$$

$$\Rightarrow \hat{\mu} = \frac{1}{m} \sum_{i=1}^m \tilde{Y}_i$$

$$\frac{\partial}{\partial \hat{\sigma}^2} \ln[L(\vec{\tilde{Y}}; \hat{\vec{x}})] = -\frac{m}{2\hat{\sigma}^2} + \frac{1}{2\hat{\sigma}^4} \sum_{i=1}^m (\tilde{Y}_i - \hat{\mu})^2 = 0$$

$$\Rightarrow \hat{\sigma}^2 = \frac{1}{m} \sum_{i=1}^m (\tilde{Y}_i - \hat{\mu})^2$$

$$E[\hat{\mu}] = E\left[\frac{1}{m} \sum_{i=1}^m \tilde{Y}_i\right] = \frac{1}{m} E\left[\sum_{i=1}^m \tilde{Y}_i\right] = \frac{1}{m} \cdot m\mu = \mu$$

$$\begin{aligned} \text{Var}[\hat{\mu}] &= \text{Var}\left[\frac{1}{m} \sum_{i=1}^m \tilde{Y}_i\right] \\ &= \frac{1}{m^2} \text{Var}\left[\sum_{i=1}^m \tilde{Y}_i\right] \\ &= \frac{1}{m^2} (m\sigma^2) = \frac{\sigma^2}{m} \end{aligned}$$

$$\text{Var}[\hat{\mu}] = E[\hat{\mu}^2] - (E[\mu])^2$$

$$E[\hat{\mu}^2] = \frac{\sigma^2}{m} + \mu$$

$$\hat{\sigma}^2 = \frac{1}{m} \sum_{i=1}^m (\tilde{Y}_i - \hat{\mu})^2 \quad \star \quad \text{sample variance}$$

$$\begin{aligned} E[\hat{\sigma}^2] &= E\left[\frac{1}{m} \sum_{i=1}^m (\tilde{Y}_i - \hat{\mu})^2\right] \\ &= \frac{1}{m} E\left[\sum_{i=1}^m (\tilde{Y}_i^2 - 2\hat{\mu}\tilde{Y}_i + \hat{\mu}^2)\right] \\ &= \frac{1}{m} E\left[\sum_{i=1}^m \tilde{Y}_i^2 - 2\hat{\mu} \sum_{i=1}^m \tilde{Y}_i + m\hat{\mu}^2\right] \\ &= \frac{1}{m} E\left[\sum_{i=1}^m \tilde{Y}_i^2 - 2m\hat{\mu}^2 + m\hat{\mu}^2\right] \\ &= \frac{1}{m} E\left[\sum_{i=1}^m \tilde{Y}_i^2 - m\hat{\mu}^2\right] \\ &= \frac{1}{m} (m \cdot E[\tilde{Y}_i^2]) - \frac{1}{m} (m \cdot E[\hat{\mu}^2]) \\ &= \mu^2 + \sigma^2 - \left(\frac{\sigma^2}{m} + \mu^2\right) \\ &= \frac{m-1}{m} \sigma^2 \Rightarrow \hat{\sigma}^2 \text{ is biased} \end{aligned}$$

$$\hat{\hat{\sigma}}^2 = \frac{1}{m-1} \sum_{i=1}^m (\tilde{Y}_i - \hat{\mu})^2 \text{ is unbiased}$$

Example :

Suppose we have data from $Y \sim B(1, p)$ (discrete random variable)

$$f(\tilde{y}_i; \tilde{x}) = x^{\tilde{y}_i} (1-x)^{1-\tilde{y}_i} \text{ where } x=p, \tilde{y}_i=0,1$$

$$= \begin{cases} x & \text{if } \tilde{y}_i=1 \\ 1-x & \text{if } \tilde{y}_i=0 \end{cases}$$

e.g. toss a coin n times.

$$L(\hat{\tilde{y}}; \tilde{x}) = \prod_{i=1}^m x^{\tilde{y}_i} (1-x)^{(1-\tilde{y}_i)}$$

$$\ln L(\tilde{y}; x) = \sum_{i=1}^m \tilde{y}_i \ln x + \sum_{i=1}^m (1-\tilde{y}_i) \ln(1-x)$$

$$J = \text{Max } \ln L(\tilde{y}; \hat{x})$$

$$\text{Necessary condition: } \frac{\partial}{\partial x} L(\tilde{y}; \hat{x}) = 0$$

$$\Rightarrow \sum_{i=1}^m \tilde{y}_i \frac{1}{x} + \sum_{i=1}^m (1-\tilde{y}_i) \frac{1}{1-x} = 0$$

$$(1-\hat{x}) \sum_{i=1}^m \tilde{y}_i + \hat{x} \sum_{i=1}^m (1-\tilde{y}_i) = 0$$

$$\sum_{i=1}^m \tilde{y}_i - \hat{x} \sum_{i=1}^m \tilde{y}_i + m\hat{x} - \hat{x} \sum_{i=1}^m \tilde{y}_i = 0$$

$$\Rightarrow \hat{x} = \frac{\sum_{i=1}^m \tilde{y}_i}{m} \quad \text{This is the estimation for } p$$

$$E[\tilde{y}_i] = p$$

$$E[\hat{x}] = E\left[\frac{1}{m} \sum_{i=1}^m \tilde{y}_i\right] = p \quad \text{This is an unbiased estimation}$$

Information matrix (scalar in this case)

$$L(\tilde{y}; x) = \prod_{i=1}^m f(\tilde{y}_i; x)$$

$$\ln L(\tilde{y}; x) = \sum_{i=1}^m \ln f(\tilde{y}_i; x)$$

$$\frac{\partial}{\partial x} \ln L(\tilde{y}; x) = \sum_{i=1}^m \frac{\partial}{\partial x} \ln f(\tilde{y}_i; x)$$

$$\frac{\partial^2}{\partial x^2} \ln L(\tilde{y}; x) = \sum_{i=1}^m \frac{\partial^2}{\partial x^2} \ln f(\tilde{y}_i; x)$$

$$F = \frac{m}{x(1-x)}$$

$$f(\tilde{y}_i; \tilde{x}) = x^{\tilde{y}_i} (1-x)^{1-\tilde{y}_i}$$

$$\ln f(\tilde{y}_i; x) = \tilde{y}_i \ln x + (1-\tilde{y}_i) \ln(1-x)$$

$$\frac{\partial}{\partial x} \ln f(\tilde{y}_i; x) = \frac{\tilde{y}_i}{x} - \frac{1-\tilde{y}_i}{1-x}$$

$$\frac{\partial^2}{\partial x^2} \ln f(\tilde{y}_i; x) = -\frac{\tilde{y}_i}{x^2} - \frac{1-\tilde{y}_i}{(1-x)^2}$$

$$E\left[\frac{\partial^2}{\partial x^2} \ln f(\tilde{y}_i; x)\right] = -\frac{1}{x} - \frac{1}{1-x}$$

$$F = -E\left[\frac{\partial^2}{\partial x^2} \ln f(\tilde{y}_i; x)\right] = \frac{1}{x} + \frac{1}{1-x} = \frac{1}{x(1-x)}$$

$$P = \text{Var}[\hat{x}] = \text{Var}\left[\frac{1}{m} \sum_{i=1}^m \tilde{y}_i\right] = \frac{1}{m^2} \cdot m \text{Var}(\tilde{y}_i) = \frac{1}{m} x(1-x)$$

$$\Rightarrow P = F^{-1}$$

2. Cramer - Rao inequality

$$\hat{\mu} = \frac{1}{m} \sum_{i=1}^m \tilde{Y}_i$$

Minimum variance estimation

$$\tilde{Y}_i = \mu + V_i \text{ where } V \sim N(0, \sigma^2)$$

$$\tilde{\mathbf{Y}} = H\mathbf{x} + \tilde{\mathbf{V}} \text{ where } \mathbf{x} = \mu, H = [1, \dots, 1]^T \in \mathbb{R}^{m \times 1}$$

$$R = \text{cov}[\tilde{\mathbf{V}}\tilde{\mathbf{V}}^T] = \begin{bmatrix} \sigma^2 & & \\ & \ddots & \\ & & \sigma^2 \end{bmatrix}$$

$$\hat{\mathbf{x}} = \hat{\mu} = (H^T R^{-1} H)^{-1} H^T R^{-1} \tilde{\mathbf{Y}}$$

$$= \frac{1}{m} \sum_{i=1}^m \tilde{Y}_i$$

$$f(\tilde{\mathbf{Y}}; \mu) = \left(\frac{1}{2\pi\sigma^2}\right)^{\frac{m}{2}} \exp\left(-\frac{\sum_{i=1}^m (\tilde{Y}_i - \mu)^2}{2\sigma^2}\right)$$

$$\ln f(\tilde{\mathbf{Y}}; \mu) = -\frac{m}{2} \ln(2\pi\sigma^2) - \frac{1}{2\sigma^2} \sum_{i=1}^m (\tilde{Y}_i - \mu)^2$$

$$\frac{\partial}{\partial \mu} \ln f(\tilde{\mathbf{Y}}; \mu) = \frac{1}{\sigma^2} \sum_{i=1}^m (\tilde{Y}_i - \mu)$$

$$\frac{\partial^2}{\partial \mu^2} \ln f(\tilde{\mathbf{Y}}; \mu) = -\frac{m}{\sigma^2}$$

Let $F = \frac{m}{\sigma^2}$ Information matrix (Fisher)

$$P = E\{(\hat{\mathbf{x}} - \bar{\mathbf{x}})(\hat{\mathbf{x}} - \bar{\mathbf{x}})^T\} \quad \text{Covariance of } \hat{\mathbf{x}}$$

$$= E\{(\hat{\mu} - \mu)^2\}$$

$$= E\{\hat{\mu}^2 - 2\mu\hat{\mu} + \mu^2\}$$

$$= E[\hat{\mu}^2] - 2\mu E[\hat{\mu}] + \mu^2$$

$$= \frac{\sigma^2}{m} + \mu^2 - 2\mu^2 + \mu^2$$

$$= \frac{\sigma^2}{m}$$

$$P = F^{-1}$$

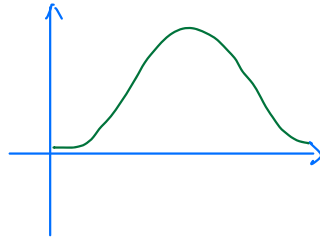
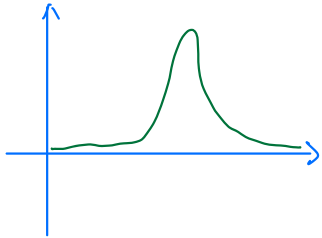
$$P \geq F^{-1}$$

Define Fisher matrix as

$$F = E \left\{ \left[\frac{\partial}{\partial x} \ln[f(\tilde{y}; \vec{x})] \right] \left[\frac{\partial}{\partial x} \ln[f(\tilde{y}; \vec{x})] \right]^T \right\}$$

can be computed by

$$F = -E \left\{ \frac{\partial^2}{\partial \vec{x} \partial \vec{x}^T} \ln[f(\tilde{y}; \vec{x})] \right\}$$



$$F = \frac{m}{\sigma^2} \quad F^{-1} = \frac{\sigma^2}{m}$$

$$\frac{\partial}{\partial x} \ln f(\tilde{y}; \vec{x}) = \frac{1}{f} \cdot \frac{\partial}{\partial x} f$$

$$\begin{aligned} F &= E \left\{ \left[\frac{\partial}{\partial x} \ln[f(\tilde{y}; \vec{x})] \right] \left[\frac{\partial}{\partial x} \ln[f(\tilde{y}; \vec{x})] \right]^T \right\} \\ &= \int_{-\infty}^{\infty} \left[\frac{\partial}{\partial x} \ln[f(\tilde{y}; \vec{x})] \right] \left[\frac{\partial}{\partial x} \ln[f(\tilde{y}; \vec{x})] \right]^T f(\tilde{y}; \vec{x}) d\tilde{y} \\ &= \int_{-\infty}^{\infty} \frac{1}{f} \frac{\partial f}{\partial x} \left(\frac{\partial f}{\partial x} \right)^T \frac{1}{f} f d\tilde{y} \\ &= \int_{-\infty}^{\infty} \frac{1}{f} \frac{\partial f}{\partial x} \left(\frac{\partial f}{\partial x} \right)^T d\tilde{y} \end{aligned}$$

$$F^{-1} \downarrow \rightarrow F \uparrow \rightarrow f \downarrow \text{ \& } \frac{\partial f}{\partial x} \uparrow .$$

CR inequality.

$$P \equiv E\{(\hat{\vec{x}} - \vec{x})(\hat{\vec{x}} - \vec{x})^T\}$$

$$F \equiv E\left\{\left[\frac{\partial}{\partial \vec{x}} \ln f(\tilde{\vec{y}}; \vec{x})\right] \left[\frac{\partial}{\partial \vec{x}} \ln f(\tilde{\vec{y}}; \vec{x})\right]^T\right\}$$

The GR inequality is $P \succeq F^{-1}$

Start the proof by

$$\textcircled{1} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \dots \int_{-\infty}^{\infty} f(\tilde{\vec{y}}; \vec{x}) d\tilde{y}_1 d\tilde{y}_2 \dots d\tilde{y}_m = 1$$

$$\int_{-\infty}^{\infty} f(\tilde{\vec{y}}; \vec{x}) d\tilde{\vec{y}} = 1$$

$\textcircled{2}$ Take partial derivative of $\textcircled{1}$ w.r.t. $\vec{x}$

$$\frac{\partial}{\partial \vec{x}} \int_{-\infty}^{\infty} f(\tilde{\vec{y}}; \vec{x}) d\tilde{\vec{y}} = \int_{-\infty}^{\infty} \frac{\partial}{\partial \vec{x}} f(\tilde{\vec{y}}; \vec{x}) d\tilde{\vec{y}} = 0$$

$\textcircled{3}$ $\hat{\vec{x}}$ is assumed to be unbiased

$$E[(\hat{\vec{x}} - \vec{x})] = \int_{-\infty}^{\infty} \underbrace{(\hat{\vec{x}} - \vec{x})}_{\text{vector}} \underbrace{f(\tilde{\vec{y}}; \vec{x})}_{\text{scalar}} d\tilde{\vec{y}} = \vec{0}$$

$\textcircled{4}$ Take partial derivative of $\textcircled{3}$ w.r.t. $\vec{x}$

$$\begin{aligned} & \frac{\partial}{\partial \vec{x}} \int_{-\infty}^{\infty} (\hat{\vec{x}} - \vec{x}) f(\tilde{\vec{y}}; \vec{x}) d\tilde{\vec{y}} \\ &= \int_{-\infty}^{\infty} (\hat{\vec{x}} - \vec{x}) \left[\frac{\partial}{\partial \vec{x}} f(\tilde{\vec{y}}; \vec{x}) \right] d\tilde{\vec{y}} - \int_{-\infty}^{\infty} f(\tilde{\vec{y}}; \vec{x}) d\tilde{\vec{y}} = 0 \\ & \int_{-\infty}^{\infty} (\hat{\vec{x}} - \vec{x}) \left[\frac{\partial}{\partial \vec{x}} f(\tilde{\vec{y}}; \vec{x}) \right]^T d\tilde{\vec{y}} = I \end{aligned}$$

$$\textcircled{5} \frac{\partial}{\partial \vec{x}} \ln f(\tilde{\vec{y}}; \vec{x}) = \frac{1}{f} \cdot \frac{\partial}{\partial \vec{x}} f$$

$$\frac{\partial}{\partial \vec{x}} f(\tilde{\vec{y}}; \vec{x}) = \left[\frac{\partial}{\partial \vec{x}} \ln f(\tilde{\vec{y}}; \vec{x}) \right] \cdot f(\tilde{\vec{y}}; \vec{x})$$

$\textcircled{6}$ plug $\textcircled{5}$ into $\textcircled{4}$

$$\int_{-\infty}^{\infty} (\hat{\vec{x}} - \vec{x}) \left[\frac{\partial}{\partial \vec{x}} \ln f(\tilde{\vec{y}}; \vec{x}) \right]^T \underbrace{f(\tilde{\vec{y}}; \vec{x})}_{\text{scalar}} d\tilde{\vec{y}} = I$$

$$\text{Let } \vec{a} \equiv f(\tilde{\vec{y}}; \vec{x})^{1/2} (\hat{\vec{x}} - \vec{x})$$

$$\vec{b} \equiv f(\tilde{\vec{y}}; \vec{x})^{1/2} \left[\frac{\partial}{\partial \vec{x}} \ln f(\tilde{\vec{y}}; \vec{x}) \right]$$

$$\int_{-\infty}^{\infty} (\vec{a} \vec{b}^T) d\tilde{\vec{y}} = I$$

$$\begin{aligned}
 P &\equiv E\{(\hat{\vec{x}} - \vec{x})(\hat{\vec{x}} - \vec{x})^T\} \\
 &= \int_{-\infty}^{\infty} (\hat{\vec{x}} - \vec{x})(\hat{\vec{x}} - \vec{x})^T f(\tilde{\vec{y}}; \vec{x}) d\tilde{\vec{y}} \\
 &= \int_{-\infty}^{\infty} \vec{a} \vec{a}^T d\tilde{\vec{y}}
 \end{aligned}$$

$$\begin{aligned}
 F &\equiv E\{[\frac{\partial}{\partial \vec{x}} \ln[f(\tilde{\vec{y}}; \vec{x})]] [\frac{\partial}{\partial \vec{x}} \ln[f(\tilde{\vec{y}}; \vec{x})]]^T\} \\
 &= \int_{-\infty}^{\infty} \vec{b} \vec{b}^T d\tilde{\vec{y}}
 \end{aligned}$$

From ⑥, we have

$$I = \int_{-\infty}^{\infty} \vec{a} \vec{b}^T d\tilde{\vec{y}}$$

For any random vectors $\vec{a}, \vec{\beta}$

$$\vec{a}^T I \vec{\beta} = \int_{-\infty}^{\infty} \vec{a}^T \vec{a} \vec{b}^T \vec{\beta} d\tilde{\vec{y}}$$

Schwartz inequality

$$[\int_{-\infty}^{\infty} g(\tilde{\vec{y}}; \vec{x}) h(\tilde{\vec{y}}; \vec{x}) d\tilde{\vec{y}}]^2 \leq \int_{-\infty}^{\infty} g^2(\tilde{\vec{y}}; \vec{x}) d\tilde{\vec{y}} \int_{-\infty}^{\infty} h^2(\tilde{\vec{y}}; \vec{x}) d\tilde{\vec{y}}$$

$$\text{Let } g(\tilde{\vec{y}}; \vec{x}) = \vec{a}^T \vec{a} \quad h(\tilde{\vec{y}}; \vec{x}) = \vec{b}^T \vec{\beta}$$

$$[\int_{-\infty}^{\infty} \vec{a}^T (\vec{a} \vec{b}^T) \vec{\beta} d\tilde{\vec{y}}]^2 \leq \int_{-\infty}^{\infty} \vec{a}^T \vec{a} \vec{a} \vec{a} d\tilde{\vec{y}} \int_{-\infty}^{\infty} \vec{\beta}^T \vec{b} \vec{b}^T \vec{\beta} d\tilde{\vec{y}}$$

$$[\vec{a}^T \vec{\beta}]^2 \leq (\vec{a}^T P \vec{a})(\vec{\beta}^T F \vec{\beta})$$

$$\text{Let } \vec{\beta} = F^{-1} \vec{a}; \quad \vec{\beta}^T = \vec{a}^T (F^{-1})^T$$

$$(\vec{a}^T P \vec{a})(\vec{a}^T (F^{-1})^T F F^{-1} \vec{a}) \geq [\vec{a}^T F^{-1} \vec{a}]^2$$

$$\vec{a}^T P \vec{a} \vec{a}^T (F^{-1})^T \vec{a} - \vec{a}^T F^{-1} \vec{a} \vec{a}^T F^{-1} \vec{a} \geq 0$$

$$[\vec{a}^T P \vec{a} - \vec{a}^T F^{-1} \vec{a}] \vec{a}^T F^{-1} \vec{a} \geq 0$$

$$\vec{a}^T P \vec{a} - \vec{a}^T F^{-1} \vec{a} \geq 0$$

$$\vec{a}^T (P - F^{-1}) \vec{a} \geq 0$$

$$P - F^{-1} \succeq 0$$

$$P \succeq F^{-1}$$

When are we gonna use this

1. Evaluate the estimators

2. feasibility studies

Example (CRLB cannot be attained)

$\{\tilde{y}_i\} \sim N(\mu, \sigma^2)$, given μ , estimate σ^2

$$F = -E \left\{ \frac{\partial^2}{\partial \vec{x} \partial \vec{x}^T} \ln[f(\vec{y}; \vec{x})] \right\}$$

$$f(\tilde{y}_i; \vec{x}) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left[-\frac{(\tilde{y}_i - \mu)^2}{2\sigma^2}\right]$$

$$\ln f(\tilde{y}_i; \vec{x}) = -\frac{1}{2} \ln(2\pi\sigma^2) - \frac{(\tilde{y}_i - \mu)^2}{2\sigma^2}$$

$$\frac{\partial}{\partial \sigma^2} \ln f(\tilde{y}_i; \vec{x}) = -\frac{1}{2} \frac{1}{\sigma^2} + \frac{(\tilde{y}_i - \mu)^2}{2(\sigma^2)^2}$$

$$\frac{\partial^2}{\partial (\sigma^2)^2} \ln f = \frac{1}{2(\sigma^2)^2} - \frac{2(\tilde{y}_i - \mu)^2}{2(\sigma^2)^3}$$

$$F = -E \left[\frac{1}{2\sigma^4} - \frac{(\tilde{y}_i - \mu)^2}{\sigma^6} \right]$$

$$= \frac{1}{2\sigma^4}$$

$$\text{The lower bound is } F^{-1} = \frac{2\sigma^4}{m}$$

The unbiased estimation is $\hat{\sigma}^2 = \frac{1}{m-1} \sum_{i=1}^m (\tilde{y}_i - \hat{\mu})^2$

$$\text{Var}(\hat{\sigma}^2) = \text{Var}\left(\frac{1}{m-1} \sum_{i=1}^m (\tilde{y}_i - \hat{\mu})^2\right)$$

$$= \frac{2\sigma^4}{m-1} = P > F^{-1}$$

3. Bayesian Estimation

In textbook 2.5 & 2.6, $P(\tilde{\mathbf{y}}|\tilde{\mathbf{x}})$ means likelihood function like $f(\hat{\mathbf{y}}; \tilde{\mathbf{x}})$ in the notes

Basic idea: consider the parameters (to be estimated) are random variables with a priori distribution
Bayesian estimation will combine the a priori information with measurements

Bayes' theorem

$$P(\tilde{\mathbf{x}}|\tilde{\mathbf{y}}) = \frac{P(\tilde{\mathbf{y}}|\tilde{\mathbf{x}})P(\tilde{\mathbf{x}})}{P(\tilde{\mathbf{y}})}$$

a posteriori distribution of $\tilde{\mathbf{x}}$

3.1 Maximum a posteriori estimation (MAP)

$$J = \max \ln P(\tilde{\mathbf{x}}|\tilde{\mathbf{y}})$$

Note $P(\tilde{\mathbf{y}})$ does not depend on $\tilde{\mathbf{x}}$

Maximize J is equivalent to

$$\begin{aligned} J_{\text{map}} &= \max \ln P(\tilde{\mathbf{y}}|\tilde{\mathbf{x}})P(\tilde{\mathbf{x}}) \\ &= \max \ln P(\tilde{\mathbf{y}}|\tilde{\mathbf{x}}) + \ln P(\tilde{\mathbf{x}}) \end{aligned}$$