

9/15/2013, Sunday

(P1)

• ~~Point estimator~~

MSE & bias-var tradeoff

example

$$MSE = E|\theta - \hat{\theta}|^2$$

$$\begin{aligned} \text{(fact: } E\cancel{X} \text{ VAR}(X) &= EX^2 - (EX)^2 \\ \Rightarrow EX^2 &= \text{VAR}(X) + (EX)^2 \text{)} \\ &= \text{VAR}(\theta - \hat{\theta}) + [E(\theta - \hat{\theta})]^2 \\ &= \text{VAR}(\hat{\theta}) + [\text{BIASE}(\hat{\theta}, \theta)]^2 \end{aligned}$$

Example biased and variance of normal estimator.

let $x_1, \dots, x_n$ iid $\sim N(\mu, \sigma^2)$

(example 1:
digital thermometer)

$$\left\{ \begin{aligned} \bar{x} &= \text{estimator for the mean} \\ &= \frac{1}{n} \sum_{i=1}^n x_i \\ S^2 &= \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 \end{aligned} \right.$$

$$\text{now: } E\bar{x} = E\left[\frac{1}{n} \sum_{i=1}^n x_i\right] = \frac{1}{n} \sum_{i=1}^n E x_i = \frac{1}{n} \sum_{i=1}^n \mu = \frac{n\mu}{n} = \mu$$

$$E S^2 = E\left[\frac{\cancel{(n-1)} \sigma^2}{n-1} \cdot \frac{(n-1) S^2}{\sigma^2}\right] = \frac{\sigma^2}{n-1} \cdot n-1 = \sigma^2$$

$$\boxed{\text{fact: } \frac{(n-1) S^2}{\sigma^2} \sim \chi^2_{n-1}} \quad \begin{aligned} \text{mean: } (n-1) \\ \text{var: } 2(n-1) \end{aligned}$$

$$\cancel{MSE(S^2)} =$$

$$\cancel{MSE(S^2)} = \cancel{E S^2} + (\cancel{\text{VAR}(S^2)} + [\cancel{\text{BIASE}(S^2, \sigma^2)}]^2)$$

$$\cancel{MSE(\bar{x})} = \cancel{E \bar{x}^2}$$

$$\text{BIASE}(\bar{x}) = E\bar{x} - \mu = 0$$

$$\text{BIASE}(S^2) = E S^2 - \sigma^2 = 0$$

$$\begin{aligned}
 \text{VAR}(\bar{x}) &= \text{VAR}\left(\frac{1}{n} \sum_{i=1}^n x_i\right) \\
 &= \frac{1}{n^2} \text{VAR}\left(\sum_{i=1}^n x_i\right) \\
 &= \frac{1}{n^2} \sum_{i=1}^n \text{VAR}(x_i) \\
 &= \frac{1}{n^2} \cdot \sum_{i=1}^n \sigma^2 = \frac{n\sigma^2}{n^2} = \boxed{\frac{\sigma^2}{n}}
 \end{aligned}$$

$$\begin{aligned}
 \text{VAR}(S^2) &= \text{VAR}\left(\frac{\cancel{\frac{1}{n-1}} \cdot \cancel{(n-1)}}{\frac{\sigma^2}{n-1} \cdot \boxed{\frac{(n-1)S^2}{\sigma^2}}}\right) \\
 &= \left[\frac{\sigma^2}{n-1}\right]^2 \cdot 2(n-1) \quad W \sim \chi^2_{n-1} \\
 &= \frac{2\sigma^4}{n-1}
 \end{aligned}$$

$$\text{MSE}(\bar{x}) = \text{BIASE}(\bar{x}) + \text{VAR}(\bar{x}) = 0 + \frac{\sigma^2}{n} = \frac{\sigma^2}{n}$$

$$\text{MSE}(S^2) = \text{BIASE}(S^2) + \text{VAR}(S^2) = 0 + \frac{2\sigma^2}{n-1} = \frac{2\sigma^2}{n-1}$$

Estimator #2

$$\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 = \frac{n-1}{n} S^2$$

$$E\hat{\sigma}^2 = E\left(\frac{n-1}{n} S^2\right) = E\left[\frac{\cancel{\frac{1}{n-1}} \cdot (n-1) S^2}{\cancel{\frac{\sigma^2}{n-1}} \cdot \frac{\sigma^2}{\sigma^2}}\right]$$

$$= \frac{\sigma^2}{n} E\left[\frac{(n-1) S^2}{\sigma^2}\right]$$

$$W \sim \chi^2_{n-1}$$

$$= \frac{\sigma^2}{n} \cdot (n-1)$$

$$= \frac{n-1}{n} \sigma^2$$

$$\text{VAR}(\hat{\sigma}^2) = \text{VAR}\left(\frac{n-1}{n} S^2\right) = \text{VAR}\left(\frac{\sigma^2}{n} \cdot \frac{(n-1) S^2}{\sigma^2}\right)$$

$$= \text{VAR}\left(\frac{n-1}{n}\right)^2 = \frac{\sigma^4}{n^2} \cdot \text{VAR}\left(\frac{(n-1) S^2}{\sigma^2}\right)$$

$$= \frac{\sigma^4}{n^2} \cdot 2(n-1) = \frac{2(n-1)\sigma^4}{n^2}$$

smaller variance

$$\frac{n-1}{n^2} < \frac{n-1}{n(n-1)} = \frac{1}{n-1}$$

$$\begin{aligned}
 \text{MSE}(\hat{\sigma}^2) &= 2\epsilon \text{BIASE}(\hat{\sigma}^2) + \text{VAR}(\hat{\sigma}^2) \\
 &= \left(E(\hat{\sigma}^2) - \sigma^2 \right)^2 + \text{VAR}(\hat{\sigma}^2) \\
 &= \left[\frac{n-1}{n} \sigma^2 - \sigma^2 \right]^2 + \frac{2(n-1)\sigma^4}{n^2} \\
 &= \left[\frac{-1}{n} \sigma^2 \right]^2 + \frac{2(n-1)\sigma^4}{n^2} \\
 &= \frac{\sigma^4}{n^2} + \frac{2(n-1)\sigma^4}{n^2} \\
 &= \frac{(2n-1)\sigma^4}{n^2} \\
 &< \left(\frac{2}{n-1} \right) \sigma^4 = \text{MSE}(S^2)
 \end{aligned}$$

$$\begin{aligned}
 \frac{2n-1}{n^2} - \frac{2}{n-1} &= \frac{(2n-1)(n-1) - 2n^2}{n^2(n-1)} \\
 &= \frac{2n^2 + 1 - 2n - n - 2n^2}{n^2(n-1)} \\
 &= \frac{1-3n}{n^2(n-1)} < 0 \quad (n > 2) \\
 &\quad n = 1, 2, \dots
 \end{aligned}$$

$$\Rightarrow \frac{2n-1}{n^2} < \frac{2}{n-1}$$

$\hat{\sigma}^2$ is a biased estimator, but
with smaller variance, than S^2
and also smaller MSE.

methods for finding estimators

(i) method of moments (MOM)

- perhaps the oldest.
- date back to Karl Pearson in late 1800
- simple to use.
- usually not the best.

$X_1, \dots, X_n$ samples from population $\sim f(x|\theta_1, \dots, \theta_k)$

- MOM estimator found by equating first k sample moments to the corresponding k "theoretical" moments calculated from pdf.

↙ a number

↘ equations involving k unknowns.

$$m_1 = \frac{1}{n} \sum_{i=1}^n X_i = EX$$

$$\frac{1}{n} \sum_{i=1}^n X_i^2 = EX^2$$

⋮

EX1 (normal) $X_1, \dots, X_n$ iid $N(\mu, \sigma^2)$

then to estimate μ and σ^2

$$\frac{1}{n} \sum_{i=1}^n X_i = \hat{\mu}$$

$$(\approx EX^2) \quad \frac{1}{n} \sum_{i=1}^n X_i^2 = \text{VAR}(X_i) + (EX_i)^2 = \hat{\sigma}^2 + \hat{\mu}^2$$

$$\Rightarrow \begin{cases} \hat{\mu} = \frac{1}{n} \sum_{i=1}^n X_i \\ \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n X_i^2 - \left(\frac{1}{n} \sum_{i=1}^n X_i \right)^2 \end{cases}$$

Ex 2 Binomial.

$X_1, \dots, X_n$ iid Binomial (k, p)

$$P(X_i = x | p) = \binom{k}{x} p^x (1-p)^{k-x}, \quad x = 0, 1, \dots, k.$$

to estimator p

$$\frac{1}{n} \sum_{i=1}^n X_i = k \hat{p} \Rightarrow \hat{p} = \frac{\frac{1}{n} \sum_{i=1}^n X_i}{k}$$

(if $n = 1$, only count once.

$$X = k \hat{p} \Rightarrow \hat{p} = \frac{X}{k}$$

Ex 3 $X_1, \dots, X_n \sim \exp(\lambda)$

$$E X = 1/\lambda = \frac{1}{n} \sum_{i=1}^n X_i \quad \lambda = 1/E \left(\frac{1}{n} \sum_{i=1}^n X_i \right)$$

② Max likelihood

• most popular for deriving estimators

• if $X_1, \dots, X_n$ iid $f(x, \theta_1, \dots, \theta_k)$

likelihood function

$$L(\theta | x) = \prod_{i=1}^n f(x_i | \theta_1, \dots, \theta_k)$$

$$\hat{\theta} = \underset{\theta \in \Theta}{\operatorname{argmax}} L(\theta | x)$$

↓
ML estimator.

(parameter point where the observed samples is most likely)

eg ~~f(x)~~ $X_1, \dots, X_n$ Bernoulli(p)

$$f(x; p) = \begin{cases} p^x (1-p)^{1-x} & , x=1,0 \\ =0 \end{cases}$$

$$\begin{aligned} L(p) &= \prod_{i=1}^n f(x_i; p) \\ &= \prod_{i=1}^n p^{x_i} (1-p)^{1-x_i} \\ &= p^{\sum_{i=1}^n x_i} (1-p)^{n - \sum_{i=1}^n x_i} \end{aligned}$$

$$\frac{\partial}{\partial p} L(p) \log L(p) = \left(\sum_{i=1}^n x_i \right) \log p + \left(n - \sum_{i=1}^n x_i \right) \log (1-p)$$

$$\frac{\partial}{\partial p} \log L(p) = \frac{1}{p} \left(\sum_{i=1}^n x_i \right) - \frac{1}{1-p} \left(n - \sum_{i=1}^n x_i \right) \stackrel{!}{=} 0$$

$$\frac{1-p}{p} = \frac{n - \sum_{i=1}^n x_i}{\sum_{i=1}^n x_i} = \frac{1 - \bar{x}}{\bar{x}}$$

$$\begin{aligned} \Rightarrow \bar{x} - p\bar{x} &= p(1 - \bar{x}) \\ &= p - p\bar{x} \end{aligned}$$

$$\begin{aligned} \Rightarrow \bar{x} &= \hat{p} \\ \hat{p} &= \frac{1}{n} \sum_{i=1}^n x_i \end{aligned}$$

e.g. 2 $x_1, \dots, x_n \sim N(\mu, \sigma^2)$

$$L(\mu, \sigma^2) = \prod_{i=1}^n f(x_i | \mu, \sigma^2) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(x_i - \mu)^2}{2\sigma^2}}$$

$$\log L(\mu, \sigma^2) = \sum_{i=1}^n \log f(x_i | \mu, \sigma^2)$$

$$= \sum_{i=1}^n \left[-\frac{1}{2} \log(2\pi\sigma^2) - \frac{(x_i - \mu)^2}{2\sigma^2} \right]$$

$$\log L(\mu, \sigma^2) = \sum_{i=1}^n \left[-\frac{1}{2} \log(2\pi\sigma^2) - \frac{(x_i - \mu)^2}{2\sigma^2} \right]$$

$$\frac{\partial}{\partial \mu} \log L(\mu, \sigma^2) = \sum_{i=1}^n -\frac{2(x_i - \mu)}{2\sigma^2} (-1)$$

$$= \sum_{i=1}^n \frac{(x_i - \mu)}{\sigma^2} = 0 \Rightarrow \hat{\mu} = \frac{1}{n} \sum_{i=1}^n x_i$$

$$\frac{\partial}{\partial \sigma^2} \log L(\mu, \sigma^2) = \sum_{i=1}^n \left[-\frac{1}{2} \cdot \frac{1}{2\pi\sigma^2} \cdot 2\pi - (-1) \frac{(x_i - \mu)^2}{2\sigma^4} \right] = 0$$

$$\Rightarrow \sum_{i=1}^n \left[-\frac{1}{2} \frac{1}{\sigma^2} + \frac{(x_i - \mu)^2}{2\sigma^4} \right] = 0$$

$$\sum_{i=1}^n \frac{(x_i - \mu)^2}{\sigma^4} = \frac{n}{\sigma^2}$$

$$\Rightarrow \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \hat{\mu})^2$$

- why Maximum likelihood estimator?

MLE property:

when n is large, and $\hat{\theta}$ is the ML estimator for θ , then

- (i) $\hat{\theta}$ is an approximately unbiased estimator

$$E(\hat{\theta}) \approx \theta$$

- (ii) $\text{VAR}(\hat{\theta})$ is approx smallest among all estimators

- (iii) $\hat{\theta}$ has approximate normal distribution

e.g. $\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$

- Drawback:

- have to find

$$\begin{aligned} \text{argmax } L(\theta) \\ = \prod_{i=1}^n f(x_i | \theta) \end{aligned}$$

not always an easy task

- method 1: set $\frac{\partial L(\theta)}{\partial \theta} = 0$

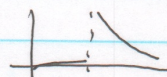
↓
Solve the equation,
not always easy

- example: $X \sim \text{uniform over } 0 \text{ to } a$

$$f(x) = 1/a, \text{ for } 0 \leq x \leq a$$

$$L(a) = \prod_{i=1}^n \frac{1}{a} = \frac{1}{a^n} \text{ for } 0 \leq x_i \leq a$$

$$\frac{dL(a)}{da} = -n a^{-(n+1)}$$



Calculus method doesn't work here
because $L(a)$ is maximized at discontinuity