

Why we like MLEs

①

- Consistency $\hat{\theta} \rightarrow^p \theta_0$
- Asymptotic normality $\sqrt{n}(\hat{\theta} - \theta_0) \rightarrow^d N(0, V_0^2)$, $V_0^2 = \{I(\theta_0)\}^{-1}$
- asymptotic efficiency $V_0^2 = \{I(\theta_0)\}^{-1}$

Consistency; we showed that the maximizer of the likelihood converges to the maximizer of the expected likelihood (Uniform conv. in prob of $\bar{\ell}$ to ℓ_0)

Asympt Normality; we used the CLT applied to

the score function and Slutsky's Theorem (obs FI $\rightarrow$ expected FI) to conclude AN of $\hat{\theta}$

Efficiency

Let's say we want to compare estimators $T^1(X^n)$ and $T^2(X^n)$

Let's say $V(T^1(X^n)) = G_{n1}^2(\theta) \xrightarrow{n \rightarrow \infty} G_1^2(\theta)$ and $V(T^2(X^n)) = G_{n2}^2(\theta) \rightarrow G_2^2(\theta)$

The asymptotic efficiency of T^1 relative to T^2 is $\frac{G_2^2(\theta)}{G_1^2(\theta)} = e_{12}(\theta)$

Note, $e_{12}(\theta)$ may be > 1 for some θ and < 1 for others!

If $g_1^2(\theta) = \frac{(g'(\theta))^2}{f_L(\theta)}$ we say that $T'(X^n)$ is

(2)

asymptotically efficient.

This typically holds for MLEs BUT we can

also get this result for other estimators

- UMVUEs (more later)

Breakdown of assumptions

(10)

Problems we can encounter ;

- Non-existence of MLE
- Multiple solutions
- Support of $p_{\theta}(x)$ depends on θ

Example of non-existence

(can happen for all x , or just some outcomes)

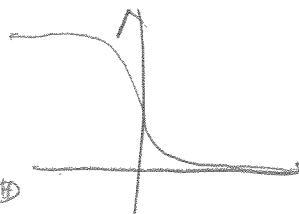
Usually when Θ not compact or $\bar{L}(\theta)$ not continuous in θ

$$X \sim \text{Be}\left(\frac{1}{1+e^{\theta}}\right)$$

$$\text{if observe } x=1 \Rightarrow L(\theta|1) = \frac{1}{1+e^{\theta}}$$

$\Rightarrow$ max occurs on boundary of Θ

$$\Rightarrow \max \hat{\theta} = -\infty !$$



Example of multiple solutions

Usually happens when θ can't be identified.

Example regression w. nonfullrank design matrix

$$y \sim N(X\theta, \Sigma) \quad \text{rank}(\Sigma) < p$$

$$(\theta_1, \dots, \theta_p)$$

Sometimes, multiple solution is a problem for small n and goes away for large enough n . (1)

Unless — # parameters increase with the sample size!

Example $\underline{X} = (X_1, \dots, X_n)$ $X_{ij} \sim N(\mu_i, \sigma^2)$ $j=1, 2$ (bivariate normal)

$n+1$ parameters $(\mu_1, \dots, \mu_n, \sigma^2) = \theta$

$$L(\theta | \underline{X}) = \prod_{i=1}^n \left(\frac{1}{\sqrt{2\pi}\sigma} \right)^2 \exp \left(-\frac{1}{2\sigma^2} \sum_{j=1}^2 (X_{ij} - \mu_i)^2 \right)$$

$$\Rightarrow \hat{\mu}_i = \frac{1}{2} (X_{i1} + X_{i2}) \rightarrow \text{no convergence here as } n \rightarrow \infty!$$

MLEs

$$\hat{\sigma}^2 = \frac{1}{2n} \sum_{i=1}^n \sum_{j=1}^2 (X_{ij} - \hat{\mu}_i)^2$$

$$\hat{\sigma}^2 \rightarrow \frac{\sigma^2}{2} ! \quad \text{since} \quad \hat{\sigma}^2 = \frac{1}{4n} \sum_{i=1}^n (X_{i1} - X_{i2})^2$$

$$= \frac{1}{4n} \sum_{i=1}^n \left(\frac{X_{i1} - X_{i2}}{\sqrt{2}\sigma} \right)^2 \cdot 2\sigma^2$$

$$\sim N(0, 1)$$

$$= \frac{\sigma^2}{2} \cdot \frac{1}{n} \sum_{i=1}^n (N(0, 1))^2$$

$$\underbrace{\chi^2_1}_{\rightarrow p=1} \text{ (mean 1)}$$

$$\rightarrow \frac{\sigma^2}{2} !$$

When $\dim(\theta)$ increases with n we can have convergence to values other than θ_0 , or no convergence at all.

Example ctd

(12)

$X^n = (X_1, \dots, X_n)$ X_i iid with $X_i = (X_{i1}, \dots, X_{in})$, $X_{ij} \sim N(\mu_i, \sigma^2)$
 $j=1, \dots, n$

Again, have $n+1$ parameters to estimate.

$$L(\theta | X^n) = \prod_{i=1}^n \left(\frac{1}{\sqrt{2\pi\sigma^2}} \right)^n \exp\left(-\frac{1}{2\sigma^2} \sum_{j=1}^n (X_{ij} - \mu_i)^2\right)$$

$$\Rightarrow \hat{\mu}_i = \frac{1}{n} \sum_{j=1}^n X_{ij} \quad i=1, \dots, n \quad \xrightarrow{P} \mu_i \text{ by WLLN}$$

$$\Rightarrow \hat{\sigma}^2 = \frac{1}{n^2} \sum_i \sum_j (X_{ij} - \hat{\mu}_i)^2 \quad \xrightarrow{P} \hat{\sigma}^2 \rightarrow \sigma^2$$

When support of $p_\theta(x)$ depends on θ (Example of breakdown of assumptions)

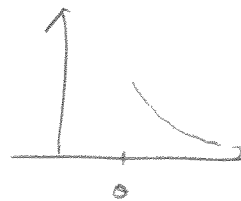
(3)

Here, we can usually get consistency, but not normality

Example

$X_i \sim$ iid shifted exponential

$$p_\theta(x) = e^{-(x-\theta)} \mathbb{I}\{x \geq \theta\}$$



$$\Rightarrow L(\theta|x) = \exp\left(-\sum_{i=1}^n (x_i - \theta)\right) \mathbb{I}\{\min x_i \geq \theta\}$$

$$\Rightarrow \text{MLE is } \min(x_i) = X_{(1)} = \hat{\theta}$$

Now, here $p_{\theta_0}(|X_{(1)} - \theta| > \varepsilon)$

$$= p_{\theta_0}(X_{(1)} > \theta_0 + \varepsilon) + p_{\theta_0}(X_{(1)} < \theta_0 - \varepsilon)$$

$$\stackrel{\text{line } n}{\text{for all } n} \prod_{i=1}^n p_{\theta_0}(X_i > \theta_0 + \varepsilon) = \exp(-n\varepsilon) \rightarrow 0$$

so $\hat{\theta}$ is consistent

However, $\sqrt{n}(X_{(1)} - \theta_0)$ cannot be normal since by definition $X_{(1)}$ is always greater than θ_0 .

Here,

$$p_{\theta_0}(n(X_{(1)} - \theta_0) \geq a) = p_{\theta_0}(X_{(1)} \geq \frac{a}{n} + \theta_0)$$

$$= \left(p_{\theta_0}(X_i \geq \frac{a}{n} + \theta_0)\right)^n = \exp(-a)$$

$$\text{so } n(X_{(1)} - \theta_0) \rightarrow^d \text{Exp}(1)$$

(faster rate of convergence here than $\sqrt{n}$)

Before we discuss what happens to MLE procedure

①

when we use the wrong model assumptions, we make a ~~debate~~ and discuss the general estimation problem.

Want to estimate $g(\theta)$ using estimator $T(X^n)$.

Is $T(X^n)$ a good estimator? "Good" means what?

First, what other type of estimators are there?

⊗ Method of Moments

Look @ r first moments $E_\theta(X^j)$ $j=1, \dots, r$

Given $p_\theta(x)$, $E_\theta(X^j) = m_j(\theta)$.

$m_j(\theta)$ can be estimated by the sample moment $\frac{1}{n} \sum_{i=1}^n X_i^j$

$\Rightarrow$ Equation system $m_j(\theta) = \bar{X}^j \quad j \geq 1$

$\Rightarrow$ solve for θ .

⊖ Note, can have many choices for estimators of θ here.

⊕ Usually easy to work with and often consistent.

⊗ Via a loss function, e.g. least squares.

$$\min_{\theta} \sum (y_i - g_i(\theta))^2$$

(Many other variants of loss fun to use, included penalized loss e.g. $\|y - g(\theta)\|^2 + \lambda \|g(\theta)\|^p$ etc.)

Comparing estimates

②

Ok, so we have an estimator $T(X^n)$ for $g(\theta)$ (estimator)

Is $T(X^n)$ a 'good' estimator?

Commonly used criterion is the $\boxed{MSE(T, \theta) = E_{\theta} (T(X^n) - g(\theta))^2}$

$\Rightarrow$ usually easy to work with in contrast with
eg. $E_{\theta} |T(X^n) - g(\theta)|$ or other norms.

Now, expanding the square we have

$$MSE(T, \theta) = \text{Var}_{\theta} T(X^n) + B^2(T, \theta)$$

$$\underbrace{\quad}_{\text{bias}^2} = (E_{\theta}(T(X^n)) - g(\theta))^2$$

Now, using $MSE(T, \theta)$ we want to find the best estimator.

Problem $MSE(T, \theta)$ usually depends on value of true θ !

So, while an estimator $T^1(X^n)$ may be better in terms of $MSE(T, \theta)$ than $T^2(X^n)$ for some θ , it may be worse for others!

No uniformly optimal estimator!

Definition: inadmissible

If two estimators S and T are such that $MSE(T, \theta) \leq MSE(S, \theta)$ with strict inequality for some θ , T is better than S and S is called inadmissible.

There are no best estimators though.

If $\theta = \theta_0$ fixed value, then $S(X) = \theta_0$ has $MSE(S, \theta_0) = 0$
and to beat it for any value $\theta_0 \Rightarrow T(X)$ must have
 $MSE(T, \theta) = 0$ everywhere $\rightarrow$ impossible task.

Problem is that we have put no restrictions on form of estimators.

$\Rightarrow$ What we tend to focus on is classes of estimators that can be compared using $MSE(\theta)$ (or some other risk function) and optimality is possible



Alt 1

Restrict to
unbiased
estimators



Alt 2

Use an average
risk as your criteria
 $\int MSE(T, \theta) \pi(\theta) d\theta$
(Bayes)



Alt 3

Use worst-case
scenario as
your criteria
 $\max_{\theta} MSE(T, \theta)$
(minimax)

Let's focus on Alt 1, unbiased estimators.

$$\text{Here } E(T(X^n)) = g(\theta) \text{ so } B(T, \theta) = 0$$

$$\text{and } \text{MSE}(T, \theta) = \text{Var}_\theta T(X^n)$$

Now, within this class of estimators we say that T is the uniformly minimum variance unbiased estimator (UMVUE)

$$\text{if } V(T(X^n)) \leq V(S(X^n)) \text{ for any unbiased estimator } S.$$

Which T is optimal? Need the following:

Rao-Blackwell Theorem

If $T(X^n)$ is sufficient for θ [Remember: means $X|T(X)$ indep of θ]
and $S(X^n)$ is an estimator for $g(\theta)$;

By conditioning on $T(X^n)$ we can always get a better, or as good as S , estimator for $g(\theta)$.

$$\text{Call this } T^*(X^n) = E(S(X^n) | T(X^n))$$

$$\text{Why? } \forall \theta \quad \text{MSE}(T^*, \theta) = E_\theta (T^*(X^n) - g(\theta))^2 \leq E_\theta (S(X^n) - g(\theta))^2$$

with strict inequality only if $T^*(X^n) \neq S(X^n)$.

$$\text{Follows from } E_\theta [E(S(X^n) | T(X^n))] = E(S(X^n)) \text{ so } B(T, \theta) = B(S, \theta) \\ = E(T^*(X^n))$$

Now $\text{Var}(E(S|T)) \leq \text{Var}(S)$ since conditioning reduces variability $\Rightarrow E(E(S|T) - E(S|T))^2 = E(E(S|T) - E(S))^2 \leq E(S - E(S))^2$ #

Completeness

T is complete if $E_0(h(T)) = 0, \forall \theta$ is only possible if $h(T) = 0$

$\Rightarrow$ Main result, due to Lehmann & Scheffé

If $T(X^n)$ is complete sufficient and $S(X^n)$ is an unbiased estimator of $g(\theta)$, then $T^*(X^n) = E(S(X^n) | T(X^n))$ is an UMVUE of $g(\theta)$. (T^* is unique if $V(T^*) < \infty \forall \theta$)

Proof

From Rao-Blackwell, we know $V(T^*) < V(S)$ for an unbiased estimator S .

Now, we need to show it didn't matter what S was so we can't improve on T^* with another S .

But from completeness, if $T_1^* = h_1(T)$ and $T_2^* = h_2(T)$ are unbiased estimators of $g(\theta)$ (from S_1 and S_2), then

$$E(h_1(T) - h_2(T)) = E(h_1(T)) - E(h_2(T)) = 0 \quad \forall \theta$$

But since we assumed T was complete $\Rightarrow h_1 = h_2$ must hold and it didn't matter what S was.

Key to the UMVUE result is having a complete sufficient statistic, $\rightarrow$ usually true for exponential families.

- Use 1: Find unbiased estimator $S(X^n)$ and derive $E(S(X^n) | T(X^n))$
- Use 2: If we have such as $T(X^n)$, and we find an estimator $h(T(X^n))$ which is unbiased for $g(\theta) \Rightarrow h(T(X^n))$ is UMVUE.

Since $E(h(T(X^n)) | T(X^n)) = h(T(X^n))$

Exponential family

6

Let's assume we have $p_\theta(x)$ has support indep of θ and is 'nice' in form of exchanging derivatives & integrals etc...

$\Rightarrow$ If $T^*(x^n)$ is an unbiased estimator for $g(\theta)$ and attains the lower bound in CR $\forall \theta \Rightarrow p_\theta$ is a one-parameter exp. family of form $p_\theta(x) = \exp\{\eta(\theta)T^*(x) - A(\theta) + L(x)\}$

$\Leftarrow$ If p_θ is a one-parameter exp. family given by $p_\theta(x) = \exp\{\eta(\theta)T(x) - A(\theta) + L(x)\}$ then (provided $\eta(\theta)$ has cb derivatives) $T(x^n)$ achieves the bound and is a UMVU for $E_\theta T(x)$

BUT $\rightarrow$ conditions may fail

(or)

an UMVU estimator for $g(\theta)$ may exist but not hit the bound (depends on $g(\cdot)$).

Some pros and cons of MLE vs UMVUE.

- large n , usually doesn't matter, but MLEs tend to be easier to compute

- MLEs are invariant in sense that if $\hat{\theta}$ is a MLE of θ then $h(\hat{\theta})$ is an MLE of $h(\theta)$.

Unbiasedness property is not always so, i.e. $\hat{\theta}$ can be unbiased for θ but $g(\hat{\theta})$ not for $g(\theta)$.

- If p_θ exp. family, both MLE and UMVUE functions of $T(x)$

Examples

Poisson (θ)

$$p_\theta(x) = \frac{e^{-\theta} \theta^x}{x!} \quad x=0,1,\dots$$

$$\text{MLE of } \theta : L(\theta | X^n) = \prod_{i=1}^n \frac{e^{-\theta} \theta^{x_i}}{x_i!} = \exp \left\{ \underbrace{n(\theta)}_{\eta(\theta)} \underbrace{\sum x_i}_{T(X^n)} - n\theta + \sum x_i \log(\theta) \right\}$$

exp family

$$\Rightarrow l(\theta | X^n) = \text{constant} - n\theta + \sum x_i \log(\theta)$$

$$\Rightarrow \frac{\partial}{\partial \theta} l(\theta | X^n) = -n + \frac{\sum x_i}{\theta} = 0 \quad \Rightarrow \hat{\theta} = \frac{\sum x_i}{n} = \bar{x}$$

$$\Rightarrow \frac{\partial^2}{\partial \theta^2} l(\theta | X^n) = -\frac{\sum x_i}{\theta^2} \Rightarrow \mathcal{F}l(\theta) = \frac{n\theta}{\theta^2} = \frac{n}{\theta}$$

$$V(\bar{x}) = \frac{V(X)}{n} = \frac{\theta}{n} = \mathcal{F}l(\theta)^{-1} \quad \text{fish bound.}$$

Now consider $g(\theta) = \psi = e^{-\theta}$

$$\log p_\psi(x) = (\log(-\log \psi)) \sum x_i + n \log \psi + \text{constant}$$

$$\frac{\partial}{\partial \psi} \log p_\psi(x) = \frac{n}{\psi} + \frac{\sum x_i}{\psi \log \psi} = 0 \Rightarrow \log \psi = -\bar{x} \Rightarrow \psi = e^{-\bar{x}} = g(\hat{\theta})$$

What about UMVU? (sample size 1)

Consider $T(X) = \begin{cases} 1 & \text{if } X=0 \\ 0 & \text{otherwise} \end{cases}$, i.e. $T(X) = 1\{X=0\}$

$E(T(X)) = E(1\{X=0\}) = e^{-\theta}$, so $T(X)$ is an unbiased estimator for $e^{-\theta}$

$$V(T(X)) = E(T(X)^2) - e^{-2\theta}$$

$$= e^{-\theta}(1 - e^{-\theta})$$

Bound $\frac{(g'(\theta))^2}{F(\theta)} = \theta e^{-2\theta}$

$$\Rightarrow V(T(X)) \geq \theta e^{-2\theta}$$

so does not hit the bound.

But $E(1\{X=0\}|X) = 1\{X=0\}$ so $1\{X=0\}$ is UMVU for $e^{-\theta}$.

Note, for mle $\hat{\theta} = e^{-X}$ is biased $\Rightarrow$ asymptotically unbiased

but $V(e^{-X}) < V(1\{X=0\})$

(Try it at home, for $n=1, \dots$)