

RECAP

①

- $L(\theta|X^n)$ captures the information in X^n about $\theta \rightarrow$ and shape of $L(\theta|X^n)$ says something about the uncertainty about θ in X^n .
As $n \rightarrow \infty$ the $L(\theta|X^n)$ becomes more and more centered around true θ .
- $L(\theta|X^n)$ is invariant to parameterizations ($\psi = g(\theta)$) but as we saw last lecture, for some parameterizations we may not attain the information bound (best possible variance)
- Our workhorse is the score $S(\theta) = \frac{\partial}{\partial \theta} \log L(\theta|X^n) = \frac{\partial}{\partial \theta} \ell(\theta|X^n)$ from which we obtain the MLE $\hat{\theta}$ through the score equation $S(\theta) = 0$
- The precision in $\hat{\theta}$ is measured by the curvature of $\ell(\theta|X^n)$ at $\hat{\theta}$
$$I(\hat{\theta}) = -\frac{\partial^2}{\partial \theta^2} \log L(\theta|X^n) \Big|_{\hat{\theta}} \quad (\text{observed information})$$
- To work out the properties of $\hat{\theta}$, $\ell(\hat{\theta}|X^n)$ vs $\ell(\theta|X^n)$ etc we use quadratic approximation near $\hat{\theta}$.

$$\underbrace{\frac{\log \ell(\theta|X^n)}{\ell(\hat{\theta}|X^n)}}_{\text{normalized log likelihood}} \approx -\frac{1}{2} I(\hat{\theta})(\theta - \hat{\theta})^2 \iff S(\theta) \approx I(\hat{\theta})(\hat{\theta} - \theta)$$

normalized
log likelihood

$$\begin{aligned} &\updownarrow \\ &I(\hat{\theta})^{-1/2} S(\theta) \approx I(\hat{\theta})^{1/2} (\hat{\theta} - \theta) \end{aligned}$$

If approx. OK we call $\ell(\theta)$ regular

Exact if p_θ is Normal-mean model.

Using The likelihood

• Likelihood intervals

$$\{ \theta : \frac{L(\theta | X^n)}{L(\hat{\theta} | X^n)} > c \} \quad , \text{i.e. } L(\theta | X^n) \text{ not too much smaller than maximum } L(\hat{\theta} | X^n)$$

What is c ?

Not easy to choose since $L(\theta | X^n)$ data dependent and we don't know sampling distribution of $\frac{L(\theta | X^n)}{L(\hat{\theta} | X^n)}$.

BUT if we trust the quasi approx then

$$\text{Wilks statistic } W = 2 \log \frac{L(\hat{\theta} | X^n)}{L(\theta | X^n)} \sim \chi^2_1 \quad \text{under null } \theta.$$

$$\Rightarrow \text{use cutoff } c = \exp \left\{ -\frac{1}{2} \chi^2_{1, (1-\alpha)} \right\}$$

(Can also use this for testing)

$$H_0: \theta = \theta_0 \quad \text{Reject if } 2 \log \frac{L(\hat{\theta} | X^n)}{L(\theta_0 | X^n)} > \chi^2_{1, (1-\alpha)}$$

• Alternative $\Rightarrow$ push approximation to $\hat{\theta}$

$$\log \frac{L(\theta | X^n)}{L(\hat{\theta} | X^n)} \approx -\frac{1}{2} I(\hat{\theta}) (\theta - \hat{\theta})^2$$

$$LI: \{ \theta : \frac{L(\theta | X^n)}{L(\hat{\theta} | X^n)} > c \} \Rightarrow CI: \left\{ \hat{\theta} \pm \underbrace{\sqrt{-2 \log c}}_{\sqrt{\chi^2_{1, (1-\alpha)}}} I(\hat{\theta})^{-1/2} \right\}$$

Wald

$\rightarrow$ quantile of normal

• When working with multiple parameter models

$\Rightarrow$ Fact $\log \frac{L(\hat{\theta}|X^n)}{L(\hat{\theta}_0|X^n)}$ can get quite small just by chance

meaning cutoff for normalized loglikelihood
should depend on # of parameters, otherwise
we will reject $H_0: \theta = \theta_0$ too often

$$\Rightarrow 2 \log \frac{L(\hat{\theta}|X^n)}{L(\theta_0|X^n)} \sim \chi_p^2 \quad \text{under null } \theta = (\theta_1, \dots, \theta_p)$$

$\underbrace{\hspace{1cm}}$
W

$\Rightarrow$ Fact Estimating more parameters always leads
to more uncertainty

Demonstrated by comparing curvature
of likelihood of one-parameter model and
curvature of the profile likelihood

θ_1 - parameter of interest

θ_2 - other parameters

$L(\theta|X^n) = L(\theta_1, \theta_2|X^n)$ full likelihood

$\tilde{L}(\theta_1|X^n) = \max_{\theta_2} L(\theta_1, \theta_2|X^n)$ profile likelihood

one-parameter problem, assuming $\theta_2 = \hat{\theta}_2$ known

$$L(\theta|X^n) \Rightarrow I(\hat{\theta}) = \begin{pmatrix} I_{11} & I_{12} \\ I_{21} & I_{22} \end{pmatrix} \Rightarrow I(\hat{\theta})^{-1} = \begin{pmatrix} I^{11} & I^{12} \\ I^{21} & I^{22} \end{pmatrix}$$

$\sim$ (can show)
 $\tilde{L}(\theta_1|X^n) \Rightarrow \tilde{I}(\hat{\theta}_1) = (I^{11})^{-1} < I_{11}$

• Inference using the likelihood

(4)

$$S(\theta|X^n) = \frac{\partial}{\partial \theta} \log L(\theta|X^n) \Rightarrow \text{properties}$$

$$\bullet E_{\theta} S(\theta|X^n) = 0$$

$$\bullet V_{\theta} S(\theta|X^n) = FI(\theta)$$

$$\text{Where } FI(\theta) = -E_{\theta} \left(\frac{\partial^2}{\partial \theta^2} \log L(\theta|X^n) \right)$$

(provided $p_{\theta}(x)$ nice enough that we can change order of derivation and integration)

→ holds for exponential families where

$$p_{\theta}(x) = \exp \{ \eta(\theta) T(x) - A(\theta) + c(x) \}$$

→ Where we also have that observed Fisher Information @ $\hat{\theta}$ is equal to expected Fisher information @ $\hat{\theta}$

$$I(\hat{\theta}) = FI(\hat{\theta})$$

and for canonical parameters where

$$p_{\theta}(x) = \exp \{ \theta T(x) - A(\theta) + c(x) \}$$

the entire functions coincide

$$I(\theta) = FI(\theta) \quad \forall \theta.$$

• Using FI

Cramér-Rao Lower Bound

For estimator $T(X^n)$ where $E_{\theta} T(X^n) = g(\theta)$, we have

$$\text{that } \text{Var}_{\theta} T(X^n) \geq \frac{g'(\theta)^2}{n I(\theta)}$$

$\Rightarrow$ bound can be attained for exponential family

(for some parameterization $g(\theta)$ that is, not all possible parameterizations)

• Other variance bounds

UMVUEs - learning MLEs and focusing on risk measure = MSE

Lehmann-Scheffé Thm

If $T(X^n)$ is a complete, sufficient statistic

$$E h(T) = 0$$

Possible only

$$\text{if } h(T) = 0$$

$X^n | T(X^n) \text{ indep of } \theta$

then an unbiased estimator $S(X^n)$ for $g(\theta)$

can be turned into a UMVUE $T^*(X^n) = E(S(X^n) | T(X^n))$

Note, $\text{Var } T^*(X^n)$ may not attain CR bound but it will under exponential families for parameterizations where

$$p_{\theta}(X^n) = \exp\{\eta(\theta) T^*(X^n) - A(\theta) + c(x)\}$$

(6)

So for some parameterization in exponential family

the UMVUE attains the bound

* Is MLE = UMVUE?

- No, but coincide sometimes. For exponential families, both estimators involve functions of $T(X)$ [$p_\theta(x) = \exp\{\eta(\theta)T(x) - A(\theta) + c(x)\}$]
- For large n , tend to coincide.
- For small n , MLE can be biased.
- MLEs tend to be easier to "find" (existence)

MLE - asymptotic

We may not understand small sample behaviour of MLEs

- or small sample behaviour may be poor

BUT for large enough n MLEs are "best"

Why?

consistency $\hat{\theta}_{MLE} \xrightarrow{P} \theta_0$

asymptotic normality

Have $(n\hat{I}(\theta))^{-1/2} S(\theta) \rightarrow^d N(0, 1)$

$\Rightarrow \sqrt{n}(\hat{\theta} - \theta) \rightarrow^d N(0, \hat{I}(\theta)^{-1})$

Asymptotic efficiency

From CLT
& Slutsky

Note, $FI(\theta)$ depends on true value of θ

→ not so practical

⇒ we tend to use $FI(\hat{\theta})$ or $I(\hat{\theta})$ for testing, CI.

(Remember, for exponential family canonical form $FI(\theta) = I(\theta) \forall \theta$

and for exponential family $FI(\hat{\theta}) = I(\hat{\theta})$.

if not, $I(\hat{\theta})$ has better properties (coverage)

Practical use

• AN of $\hat{\theta} \Rightarrow$ inference on $\theta \rightarrow$ CI
 $\rightarrow$ testing

• LRT $\Rightarrow$ inference on nested models

$$W = 2 \log \frac{L(\hat{\theta} | X^1)}{L(\hat{\theta}_{H_0} | X^0)} \sim \chi^2_{\Delta_p} \text{ under } H_0$$

max w. some
parameter restrictions

Where Δ_p is the number
of restricted parameters
under H_0

$$H_0: \theta_{H_0}, \text{ e.g. } \begin{cases} \theta_1 = \theta_{1,0} \\ \theta_2 = ? \end{cases}$$

(proof, from AN of $\hat{\theta}$ and conditional
distribution of $\hat{\theta}_2(\theta_1 = \theta_{1,0})$