

Review 1

①

Focus on likelihood based inference

$$L(\theta | x^n) = \prod_{i=1}^n p_{\theta}(x_i) \quad \left[\text{can generalize to } p_{\theta}(x; \mathcal{R}(x;)) \right]$$

: function over θ for fixed observation x^n

Shape of $L(\theta | x^n)$ captures information about θ in this particular x^n and uncertainty about θ in this x^n

$\Rightarrow$ can try to convert to direct probabilistic statements or (w/ large-sample approximations

(empirical - frequentist methods)

- often a loss $L(\hat{\theta}(x^n) - \theta)$ 'averaged' out over repeated sampling from $F_{\theta} \Rightarrow \mathbb{E}_{F_{\theta}}[L(\hat{\theta}(x^n), \theta)]$
Risk

- discuss / base inference on long-run properties.

- 1% cover from θ 1-2% of F_{θ} -samples

- Bayesian methods

- prior $\pi(\theta) \Rightarrow$ posterior $\pi(\theta | x^n) \propto \pi(\theta) L(\theta | x^n)$

- now get probabilistic statements: $\text{CI}(x^n)$ covers θ w. prob 1- α .

• Big + w/ likelihood based inference $\rightarrow$ invariant to reparameterizations $\eta = g(\theta) \Rightarrow \hat{\eta} = g(\hat{\theta})$

• Also, to compare $L(\theta | x^n)$ for different values of $\theta \Rightarrow LR$ is invariant to reparameterizations and variable transformations,

$L(\theta|x^n)$ can have a 'complex' shape with multiple local maxima, skewness etc.

However, we are only really interested in θ 's near maximum of $L(\theta)$, or $l(\theta)$ itself near maximum (since other θ 's not supported by data).

Focus on shape of $L(\theta|x^n)$ near $\hat{\theta}_n = \underset{\theta}{\operatorname{argmax}} L(\theta|x^n)$

⊗ If $\log L(\theta|x^n) \approx l(\theta|x^n)$ is near quadratic near $\hat{\theta}$

we call $l(\theta|x^n)$ regular

→ lots of inference used in practice based on this assumption!

Reality check — plot $l(\theta|x^n)$ and its quadratic approximation and compare.

⊗ Big + as $n \rightarrow \infty$ most $l(\theta|x^n)$ look near quadratic near $\hat{\theta}_n$

The concepts we have used for derivation of tests etc.

⇒ Score, information

• $l(\theta|x^n) = \log(\text{likelihood}) = l(\theta)$ (but remember dependence on x^n)

• $S(\theta) = \frac{\partial}{\partial \theta} l(\theta|x^n) = \underline{\text{score}}$

• Definition: $\hat{\theta}_n = \hat{\theta}_{MLE} = \underset{\theta}{\operatorname{argmax}} l(\theta|x^n) = \{ \theta : S(\theta) = 0 \}$

- Curvature of $l(\theta|x^n)$

$$J(\theta) = -\frac{\partial^2}{\partial \theta^2} l(\theta)$$

- $J(\hat{\theta})$ = curvature at $\hat{\theta}$ = positive quantity by def of $\hat{\theta}$

= observed Fisher information

Compare with expected Fisher information

$$E_{\theta} \left[-\frac{\partial^2}{\partial \theta^2} l(\theta|x) \right]$$

- Quadratic approximation near $\hat{\theta}$

$$l(\theta) \approx l(\hat{\theta}) + \underbrace{S(\hat{\theta})}_{=0} (\theta - \hat{\theta}) - \frac{1}{2} J(\hat{\theta}) (\theta - \hat{\theta})^2$$

$$\Rightarrow \boxed{\underbrace{\log \frac{L(\theta)}{L(\hat{\theta})}}_{\text{normalized log likelihood}} \approx \underbrace{-\frac{1}{2} J(\hat{\theta}) (\theta - \hat{\theta})^2}_{\text{quadratic approximation}}}$$

- Derivative of LH $\Rightarrow$ $\boxed{S(\theta) \approx J(\hat{\theta}) (\hat{\theta} - \theta)}$

This is important

$\hat{\theta}$ inherits all its properties from $S(\theta)$!

So if we understand $S(\theta)$ behaviour

$\Rightarrow$ understand $\hat{\theta}$.

Looking at the score in more detail

• $S(\theta) = \frac{\partial}{\partial \theta} \log L(\theta)$, $S(\hat{\theta}) = 0$ by definition of $\hat{\theta}$

• Take expectation of $S(\theta)$ under the model p_θ

$$E_\theta[S(\theta)] = E_\theta \left[\frac{\partial}{\partial \theta} \log L(\theta | X) \right] \quad (\text{can focus on } \log L(\theta | X) \text{ since } S(\theta) \text{ is additive in } X, \text{ iid } p_\theta)$$

$$= \int \frac{\partial}{\partial \theta} \log L(\theta | x) p_\theta(x) dx$$

$$= \int \frac{\frac{\partial}{\partial \theta} L(\theta | x)}{L(\theta | x)} p_\theta(x) dx = \int \frac{\partial}{\partial \theta} p_\theta(x) dx = \frac{\partial}{\partial \theta} \int p_\theta(x) dx = 0$$

[for nice p_θ where $\frac{\partial}{\partial \theta}$ and $\int$ can change order]

• $\text{Var}_\theta S(\theta) = E_\theta[S(\theta)^2] = \int \left(\frac{\partial \log L(\theta)}{\partial \theta} \right)^2 p_\theta(x) dx$

$$= \int \left(\frac{\partial \log L(\theta | x)}{\partial \theta} \right) \frac{\frac{\partial}{\partial \theta} L(\theta | x)}{L(\theta | x)} p_\theta(x) dx = \int \left(\frac{\partial \log L(\theta | x)}{\partial \theta} \right) \left(\frac{\partial}{\partial \theta} L(\theta | x) \right) dx$$

Integration by parts

$$\frac{\partial \log L(\theta)}{\partial \theta} \left(\int \frac{\partial}{\partial \theta} p_\theta(x) dx \right) - \int \frac{\partial^2 \log L(\theta | x)}{\partial \theta^2} p_\theta(x) dx = -E_\theta \left[\frac{\partial^2 \log L(\theta | x)}{\partial \theta^2} \right]$$

$$= 0$$

$$= \mathcal{F}_T(\theta)$$

expected Fisher Information

(5)

Now, we have

$$S(\theta) \approx I(\hat{\theta})(\hat{\theta} - \theta) \quad (\text{quadratic approx of } l(\theta))$$

$$\Rightarrow I(\hat{\theta})^{-1/2} S(\theta) \approx I(\hat{\theta})^{1/2} (\hat{\theta} - \theta)$$

sampling distribution
of this?

$$S(\theta) = \sum_i \frac{\partial}{\partial \theta} \log p_{\theta}(x_i) = \sum_i y_i \quad \text{where } y_i = \frac{\partial}{\partial \theta} \log p_{\theta}(x_i)$$

- the y_i 's are iid

$$\text{- we know } E_{\theta}(y_i) = 0 \quad \text{and} \quad \text{Var}_{\theta}(y_i) = FI(\theta) = E_{\theta} \left[-\frac{\partial^2}{\partial \theta^2} \log L(\theta | X) \right]$$

$$\Rightarrow E_{\theta}(S(\theta)) = 0, \quad \text{Var}_{\theta}(S(\theta)) = n FI(\theta)$$

• Central Limit Theorem applied to $\sum y_i = S(\theta)$

$$\Rightarrow \frac{1}{\sqrt{n}} S(\theta) \rightarrow N(0, FI(\theta))$$

$$\text{or } \boxed{(n FI(\theta))^{-1/2} S(\theta) \rightarrow N(0, 1)}$$

↓

$$\text{but we have } (I(\hat{\theta}))^{-1/2} S(\theta) \approx I(\hat{\theta})^{1/2} (\hat{\theta} - \theta)$$

so...

$$\text{but } FI(\theta) = E_{\theta} \left[-\frac{\partial^2}{\partial \theta^2} \log L(\theta | X) \right] \quad \text{and} \quad I(\hat{\theta}) = -\frac{\partial^2}{\partial \theta^2} \log L(\theta | X) \Big|_{\hat{\theta}}$$

(can we just replace $I(\hat{\theta})$ by $FI(\theta)$?)

(6)

WLLN $\frac{1}{n} \sum_{i=1}^n \frac{\partial^2}{\partial \theta^2} \log f_{\theta}(X_i) \xrightarrow{P} E_{\theta} \left[\frac{\partial^2}{\partial \theta^2} \log f_{\theta}(X) \right]$

$$J_{\theta} \frac{\bar{I}(\theta)}{n} \xrightarrow{P} \bar{F} I(\theta) \quad \forall \theta$$

Slutsky's Theorem If $X_n \xrightarrow{d} X$ and $A_n \xrightarrow{P} a$

$$\text{Then } A_n X_n \xrightarrow{d} a X$$

Here $I(\hat{\theta})^{-1/2} S(\theta) = \underbrace{(n F I(\theta))^{-1/2} S(\theta)}_{\xrightarrow{d} N(0,1)} \underbrace{\left(\frac{I(\theta)}{n} \right)^{-1/2} (F I(\theta))^{1/2}}_{\xrightarrow{P} 1}$

$$\Rightarrow \text{We get } \sqrt{n}(\hat{\theta} - \theta) \rightarrow N(0, S I(\theta)^{-1})$$



$$(\hat{\theta} - \theta) \rightarrow N(0, \bar{S}(\hat{\theta})^{-1})$$

← used in practice
since $I(\theta)$ computed from data

Inference $CI: \left[\hat{\theta} \pm Z_{1-\alpha/2} \sqrt{\bar{I}(\hat{\theta})^{-1}} \right]$

Other types of interval estimates used in likelihood

(7)

$$\left\{ \theta : \frac{L(\theta)}{L(\hat{\theta})} > c \right\} \quad \text{is } \theta \text{ within reasonable likelihood range from the maximum } L(\hat{\theta})$$

What should c be?

Tricky! Data dependent $L(\theta)$ and dimension dependency makes 'standard' c difficult to find.

From quadratic approximation $\hat{\theta} \sim AN \Rightarrow \hat{\theta}^2 \sim \chi^2$

$$\Rightarrow W = 2 \log \frac{L(\hat{\theta})}{L(\theta)} \sim \chi^2, \quad \text{under null } H_0: \theta$$

(This is an approximation based on normal model)

$$\Rightarrow \left\{ \theta : 2 \log \frac{L(\hat{\theta})}{L(\theta)} < -2 \log c \right\} \Rightarrow c \text{ should be } e^{-\frac{1}{2} \chi^2_{1-\alpha}} \text{ for } 1-\alpha \text{ coverage}$$

$$\text{Both CI} = \left\{ \theta : 2 \log \frac{L(\hat{\theta})}{L(\theta)} > \chi^2_{1-\alpha} \right\} \quad (\text{Wilks})$$

and

$$CI = \left\{ \hat{\theta} \pm Z_{1-\alpha/2} I(\hat{\theta})^{-1/2} \right\}$$

are based on quadratic approximation (Wald)

In practice, Wilks works much better - allows for asymmetric CI

Multiparameter case

(8)

$$\log L(\theta) = \log L(\hat{\theta}) + S(\hat{\theta})' (\theta - \hat{\theta}) - \frac{1}{2} (\theta - \hat{\theta})' \underbrace{I(\hat{\theta})}_{p \times p \text{ matrix}} (\theta - \hat{\theta})$$

$\begin{matrix} 1 \times 1 & & 1 \times 1 & & 1 \times p & & p \times 1 & & 1 \times p & & p \times 1 \end{matrix}$
 $\theta = (\theta_1, \dots, \theta_p)$

$$I(\hat{\theta})_{ij} = \frac{\partial^2}{\partial \theta_i \partial \theta_j} \log L(\theta | X^n) \Big|_{\hat{\theta}}$$

Solve $\{\theta: S(\theta) = 0\}$ now a multiparameter problem — can be difficult

• Let's say we found solutions $\hat{\theta}$
— inference?

Wilks $W = 2 \log \frac{L(\hat{\theta})}{L(\theta)} \sim \chi^2_p$ now (approx)

$$\Rightarrow CI = \left\{ \theta: \frac{L(\theta)}{L(\hat{\theta})} > \exp\left(-\frac{1}{2} \chi^2_p(1-\alpha)\right) \right\}$$

but we rarely construct confidence regions like this — more often look at marginal sampling distribution for one $\hat{\theta}_j$ at a time

Testing an hypothesis $H_0: \theta = \theta_0$ is done

using Wilks' theorem (see GLM later)

Exponential families, UMVUE, MLE

9

Data $x^n \rightarrow$ reduce to a summary statistic $T(x^n)$
 $\rightarrow$ often of a form $\bar{x}$ or $\bar{x}^2$ etc

$T(x^n)$ is sufficient for θ if the distribution $p_\theta(x^n)$

factorizes as $p_\theta(x^n) = g(T(x^n), \theta) h(x^n)$

$\underbrace{\hspace{1cm}}$ depends on x^n only through $T(x^n)$
 $\underbrace{\hspace{1cm}}$ does not depend on θ

Exponential families

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

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

$$\text{where } T(x^n) = \sum_{i=1}^n T(x_i), \quad A_n(\theta) = nA(\theta)$$

$\underbrace{\hspace{1cm}}$ this is a nice property - easy to work with

$\eta(\theta) = \theta$ - canonical form

$$p_\theta(x^n) = \exp\{\theta T(x^n) - nA(\theta) + c(x^n)\}$$

$$\text{(can show that)} \begin{cases} E(T(x)) = A'(\theta) \\ V(T(x)) = A''(\theta) \end{cases}$$

Mean-variance relationship

Multiparameter model

$$p_{\underline{\theta}}(x) = \exp \left\{ \sum_{j=1}^p \eta_j(\underline{\theta}) T_j(x) - A(\underline{\theta}) + c(x) \right\}$$

$\underbrace{\hspace{10em}}$
function of $\underline{\theta}$
 p sufficient statistics T_j

For exponential families

$$\rightarrow l(\underline{\theta}|x) = \eta(\underline{\theta})T(x) - A(\underline{\theta}) + c(x)$$

$$S(\underline{\theta}) = \eta'(\underline{\theta})T(x) - A'(\underline{\theta}) \quad \left[S(\underline{\theta}) = T(x) - A'(\underline{\theta}) \text{ for canonical} \right]$$

$$\rightarrow \text{MLE} \Rightarrow \{ \underline{\theta} : \eta'(\underline{\theta})T(x) = A'(\underline{\theta}) \}$$

$$\rightarrow \text{Covariance } J(\underline{\theta}) = -\eta''(\underline{\theta})T(x) + A''(\underline{\theta}) \quad \left[J(\underline{\theta}) = A''(\underline{\theta}) \text{ canonical} \right]$$

For exponential families then

$$E_{\underline{\theta}}(S(\underline{\theta})) = \eta'(\underline{\theta}) \underline{E_{\underline{\theta}} T(x)} - A'(\underline{\theta}) = 0$$

$$FI(\underline{\theta}) = A''(\underline{\theta}) - \eta''(\underline{\theta}) \underline{E_{\underline{\theta}} T(x)}$$

$$* \text{ At } \hat{\underline{\theta}} : I(\hat{\underline{\theta}}) = FI(\hat{\underline{\theta}})$$

• For canonical model $J(\underline{\theta}) = FI(\underline{\theta})$ for all $\underline{\theta}$!

(11)

So, observed information = expected information for canonical families so uncertainty about θ in this sample ~ uncertainty of long-run frequency quantity

Cramér-Rao lower bound

For 'nice' distributions (like exponential p_θ)

and 'nice' estimators $T(x) = \hat{\theta}(x)$ ($E_\theta(T(x)) < \infty$ etc)

If $T(x)$ is an unbiased estimator of $g(\theta)$ with finite variance

$$\boxed{\text{Var}_\theta(T(x)) \geq g'(\theta)' F^{-1}(\theta) g'(\theta)}$$

The proof was done by using Conditional Variance = reduction in variance

$\Rightarrow$ bound can only be met by $T(x)$

that is linearly related to the score $S(\theta)$

$\Rightarrow$ exponential family

$\Rightarrow$ class of UMVUEs $\Rightarrow$ alternative to MLE

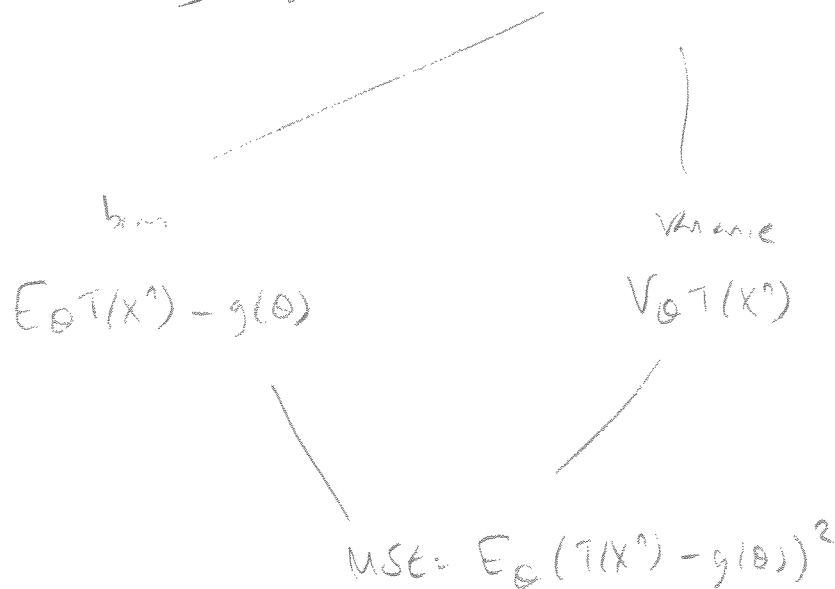
(compare in - bit)

General estimation

(13)

- parameter of interest $g(\theta)$
- estimator $T(X^n)$

Is $T(X^n)$ any good? What is "good"?



(common comb
criterion

$$MSE(T, \theta)$$

usually depends on the value of θ

so there is no single best estimator
for all $\theta \Rightarrow$ but some restrictions
on form of $T(X^n)$

(common restriction $\Rightarrow$ Unbiasedness $\rightarrow$ then look for minimum variance

$\Rightarrow$ Worst-case $\rightarrow$ min-max $\Rightarrow \min_T \max_{\theta} MSE(T, \theta)$

if stay with unbiased

(14)

→ Rao-Blackwell

if $T(X^n)$ sufficient for θ and $S(X^n)$ is an estimator of $g(\theta)$,

the new estimator $E(S(X^n) | T(X^n)) = T^*(X^n)$

has property

$$\forall \theta \quad \text{MSE}(T^*, \theta) \leq \text{MSE}(S, \theta)$$

→ Lehmann-Scheffé

if $T(X^n)$ is a complete sufficient stat for θ , and

$S(X^n)$ is an unbiased estimator of $g(\theta)$,

then $T^*(X^n)$ is an UMVUE of $g(\theta)$

(Complete: $E_\theta h(T) = 0 \quad \forall \theta$ implies $h(T) = 0$)

Construction of UMVUE

if have complete sufficient T and any unbiased S

$$\Rightarrow T^* = E(S | T) \neq$$

that part is proving T is complete — but usually rare for exponential families!

For exponential families

(15)

→ (can under one parameterization find a $T^*(X)$)

that is UMVUE and meets the CR bound!

$$\left(p_{\theta}(x) = p_{\eta}(x) \left\{ \eta(\theta) T^*(x) - A(\theta) + C(x) \right\} \right)$$

→ doesn't need work for all parameterizations, since

we may have $\hat{\theta}$ unbiased for θ but $g(\hat{\theta})$ not

unbiased for $g(\theta)$

MLE vs. UMVUE

- MLE - easy to work with
- UMVUE - requires sometimes work to find

- MLE - invariant to parameterization
- UMVUE - may find unbiased $\hat{\theta}$ but not $g(\hat{\theta})$

• For large n , MLE and UMVUE tend to agree

• For exponential families, both MLE & UMVUEs tend to be functions of the sufficient statistics $T(x)$