

Information in data — captured by the likelihood

①

→ but in what format?

Data x reduced to summary statistics $T(x)$

sufficient for estimating parameter θ : $\hat{\theta} = f(T(x))$

Definition A statistic $T(x)$ is sufficient for θ if
the distribution $X | T=t$ does not depend on θ

↓ That is, knowing $T(x)$ captures all of
the information in x about θ .

Factorization Theorem

If T is sufficient for θ we can write
the model as

$$p_{\theta}(x) = g(T(x), \theta) h(x)$$

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

Example from normal-mean model

$$\begin{aligned} X_i \sim \text{iid } N(\mu, \sigma^2) \quad \Rightarrow \quad p_{\theta}(x) &= (2\pi\sigma^2)^{-n/2} \exp \left\{ -\frac{1}{2\sigma^2} \sum (x_i - \mu)^2 \right\} \\ \theta &= (\mu, \sigma^2) \\ &= (2\pi\sigma^2)^{-n/2} \exp \left\{ -\frac{\sum x_i^2}{2\sigma^2} + \frac{\mu \sum x_i}{\sigma^2} - \frac{n\mu^2}{2\sigma^2} \right\} \end{aligned}$$

Sufficient statistics $\sum x_i^2$ and $\sum x_i$

But ... $\sum_{i=1}^m x_i$ and $\sum_{i=m+1}^n x_i$ (partial sums) are also sufficient statistics

$\Rightarrow$ concept of minimal sufficiency : maximum data reduction that still retains all the information about θ .

Definition $T(\underline{x})$ is minimally sufficient if it is a function of any other sufficient statistic.

May not be unique : any one-to-one map of minimal sufficient T is minimal sufficient

η $T(\underline{x})$ is sufficient for $\theta \Rightarrow$ likelihood based on $\underline{x} \equiv$ likelihood based on $T(\underline{x})$ (follows from Factorization Thm)

$\Rightarrow L(\theta)$ is a function of any T

$\Rightarrow$ $L(\theta)$ is minimal sufficient

Exponential Families

(2)

Important class of distributions for development of theory of statistics.

Has nice properties (but) also covers many important models used in practice.

One-parameter model

$$p_{\theta}(x) = \exp \left\{ \underbrace{\eta(\theta)}_{\text{function of } \theta} \underbrace{T(x)}_{\text{sufficient statistic}} - \underbrace{A(\theta)}_{\text{mean-function}} + \underbrace{C(x)}_{\text{normalizing constant}} \right\}$$

n iid samples

$$p_{\theta}(\underline{x}) = \exp \left\{ \eta(\theta) \underbrace{\sum_{i=1}^n T(x_i)}_{\text{nice property of exponential family}} - \underbrace{n A(\theta)}_{\tilde{A}(\theta)} + C(\underline{x}) \right\}$$

When $\eta(\theta) = \theta$ we say we have the distribution in its canonical form and it looks especially simple

$$p_{\theta}(\underline{x}) = \exp \left\{ \theta T(\underline{x}) - \tilde{A}(\theta) + C(\underline{x}) \right\}$$

Some properties

$$\Rightarrow \begin{cases} E(T(\underline{x})) = \tilde{A}'(\theta) \\ V(T(\underline{x})) = \tilde{A}''(\theta) \end{cases}$$

(proof via moment
generating functions)
(google)

This mean-variance relationship
important for selecting appropriate
model to use.

Example Normal-mean model (σ^2 known)

$$T(\underline{x}) = \bar{x}$$

$$E(\bar{x}) = \mu, \quad V(\bar{x}) = \frac{\sigma^2}{n} \text{ does not depend on } \mu$$

Poisson

$$T(\underline{x}) = \bar{x}$$

$$E(\bar{x}) = \mu, \quad V(\bar{x}) = \frac{\mu}{n} \text{ does depend on mean } \mu.$$

Multi parameter case

$$\underline{\theta} = \{\theta_1, \dots, \theta_p\}$$

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

vector of sufficient statistics $\{T_1(\underline{x}), \dots, T_p(\underline{x})\}$

vector of functions of $\underline{\theta}$ - $\eta_j(\underline{\theta})$ $j=1, \dots, p$ where
each component may involve more than one component
of $\underline{\theta}$.

Example normal mean model

5

$$p_{\theta}(x) = \exp \left\{ \frac{-x^2}{2\sigma^2} + \frac{\mu}{\sigma^2} x - \frac{1}{2} \left(\frac{\mu^2}{\sigma^2} + \ln(2\pi\sigma^2) \right) \right\}$$

$\theta = (\mu, \sigma^2)$

$$A(\theta)$$

$$\eta(\theta) = \left\{ \frac{\mu}{\sigma^2}, \frac{-1}{2\sigma^2} \right\}$$

$$T(x) = \{ x, x^2 \}$$

Likelihood & exponential families

$$L(\theta|x) = \exp \{ \eta(\theta)T(x) - A(\theta) + C(x) \}$$

~~8/2/20~~ $\log L(\theta|x) = \eta(\theta)T(x) - A(\theta) + C(x)$



score $S(\theta) = \eta'(\theta)T(x) - A'(\theta)$

$\eta(\theta) = \theta$, $S(\theta) = T(x) - A'(\theta)$

MLE $\Rightarrow \{ \theta : S(\theta) = 0 \} \Rightarrow$ solution to $\eta'(\theta)T(x) = A'(\theta)$

or if canonical form

$$T(x) = A'(\theta)$$

Curvature, Fisher information

$$I(\theta) = -\frac{\partial^2 \log L(\theta)}{\partial \theta^2} = -\eta''(\theta)T(x) + A''(\theta)$$

$\Rightarrow$ canonical form $I(\theta) = A''(\theta)$

Remember for sample size n

$$\Rightarrow T(\underline{x}) = \sum T(x_i) \quad \text{and} \quad \tilde{A}(\theta) = n A(\theta)$$

$$\Rightarrow I_n(\theta) = n A''(\theta)$$

Also remember we have

$$\text{Var}(\hat{\theta}) = I_n(\theta)^{-1} = n^{-1} (I(\theta)^{-1})$$

sample size
dependency
of variance

constant in variance
expression that depends
on shape of $L(\theta)$.

(specific family model)

Looking at the score in more detail

$$S(\theta) = \frac{\partial}{\partial \theta} \log L(\theta)$$



Take expectation
under true model
 p_θ

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

$$\stackrel{\text{Chain rule}}{=} \int \frac{\frac{\partial}{\partial \theta} L(\theta|x)}{L(\theta|x)} p_\theta(x) dx$$

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

true pr
exponential
families

for nice
distributions
(on more $\frac{\partial}{\partial \theta}$
outside $\int$)

$$\frac{\partial}{\partial \theta} \underbrace{\int L(\theta|x) dx}_{=1} = 0$$

since $\int L(\theta|x) dx = 1$

(7)

So, the expected value of $S(\theta)$ is 0

$\Rightarrow$ logical to use $\hat{\theta} = \{\theta : S(\theta) = 0\}$ as an estimate
 since we hope $S(\theta) = \frac{\partial}{\partial \theta} \log L(\theta|x)$ observed (for
 a particular x) not too far from expected

How much can $S(\theta)$ vary?

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

since
 $E_{\theta}(S(\theta)) = 0$

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

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

Integration
 $\stackrel{=}{=}$
 by parts

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

$= 0$

$$\Rightarrow \text{Var}_{\theta} S(\theta) = - \int \frac{\partial^2 \log L(\theta|x)}{\partial \theta^2} p_{\theta}(x) dx \stackrel{\text{definition}}{=} FI(\theta)$$

the expected
 Fisher information.

So, estimate $\hat{\theta} = \{ \theta : S(\theta) = 0 \}$

$\Rightarrow \hat{\theta}$ inherits properties of $S(\theta)$

— mean 0

— variance $FI(\theta)$

$\Downarrow$ mean, variance for
 θ

(more later)
 $\hat{\theta} \sim N(0, FI(0)^{-1})$

Example normal-mean model

$$S(\theta) = \frac{n}{\sigma^2} (\bar{x} - \theta) \quad (\sigma^2 \text{ known})$$

$$\Downarrow$$
$$\Rightarrow \begin{cases} E(S(\theta)) = 0 \\ V(S(\theta)) = \frac{n}{\sigma^2} \end{cases}$$

$$\text{and } L(\theta) = \frac{d}{d\theta} S(\theta) = \frac{n}{\sigma^2} \Rightarrow FI(\theta) = \frac{n}{\sigma^2}$$

Here, $L(\theta) = FI(\theta)$ happens for exponential families
in canonical form.

Exponential families

$$S(\theta) = \eta'(\theta)T(x) - A'(\theta)$$

$$\Rightarrow \hat{\theta} = \{ \theta : \eta'(\theta)T(x) = A'(\theta) \}$$

and we have

$$E(S(\theta)) = \eta'(\theta)E(T(x)) - A'(\theta) = 0$$

(9)

$$\text{Now, } I(\theta) = -\frac{\partial^2}{\partial \theta^2} \log L(\theta) = A''(\theta) - \eta''(\theta) T(x)$$

$$\text{and } FI(\theta) = E(\quad) = A''(\theta) - \eta''(\theta) E(T(x))$$

At solution to score-equation $\hat{\theta}$ (function of data $\hat{\theta} = \hat{\theta}(T(x))$)

$$I(\hat{\theta}) = A''(\hat{\theta}) - \eta''(\hat{\theta}) T(x) \quad \text{where we know } \eta'(\hat{\theta}) T(x) = A'(\hat{\theta})$$

and

$$FI(\hat{\theta}) = A''(\hat{\theta}) - \eta''(\hat{\theta}) E(T(x)) \quad \text{where we know } \eta'(\theta) E(T(x)) = A'(\theta) \text{ for all } \theta.$$

function FI
evaluated
at $\hat{\theta}$
value

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

so observed Fisher information
coincides with the expected
evaluated at the value $\hat{\theta}$

If we use canonical form we get an even
stronger result

$$I(\theta) = A''(\theta) \quad \text{so } I(\theta) = FI(\theta) \quad \text{functions coincide}$$

$$FI(\theta) = A''(\theta)$$

Interesting result. $I(\theta)$ - curvature obtained from actual sample x

→ uncertainty about estimating θ in this sample

$FI(\theta)$ is a long-run frequency property about how hard
estimation of θ is.

Turn out $FI(\theta)$ has another important meaning.

Cramér-Rao lower bound

• Let's say θ estimator $T(x)$ is unbiased
such that $E_{\theta} T(x) = \theta$

• Then $V_{\theta}(T(x)) \geq \frac{1}{FI(\theta)}$

That is, we can never do better than $FI(\theta)^{-1}$
in terms of estimation efficiency

We can meet this bound for exponential
family models!