

sufficiency (ct'd).

How to find sufficient stats?

Thm (factorization).

$\{P_\theta : \theta \in \Theta\}$, assuming there $P_\theta = \frac{dP_\theta}{d\mu}$ exists ($\forall \theta$)
(for some base msr μ).

then T is sufficient for the class

$$\exists g_\theta, h \text{ s.t. } P_\theta(x) = \underbrace{g_\theta(T(x))}_{\text{generating } T(x) \text{ based on } \theta} \cdot \underbrace{h(x)}_{\text{Conditionally on } T, \text{ generate } X \text{ indep of } \theta}.$$

(Proof assumes $\exists$ joint density).

Proof. " $\Leftarrow$ ".

$$P_\theta(x | T(x)=t) = \frac{\cancel{g_\theta(t)} h(x) \mathbb{1}_{\{T(x)=t\}}}{\int_{T(z)=t} \cancel{g_\theta(t)} h(z) d\mu(z)} = \frac{h(x) \mathbb{1}_{\{T(x)=t\}}}{\int_{T(z)=t} h(z) d\mu(z)}$$

(indep of θ).

" $\Rightarrow$ "

Construction:

marginal density

$$g_\theta(t) = p_\theta(T(X)=t)$$

$$h(x) = p_{\theta_0}(X=x | T(X)=t) \quad \text{conditional density.}$$

(for arbitrary $\theta_0 \in \Theta$).

$$p_\theta(x) = p_\theta(T(X)=t, X=x) = p_\theta(T(X)=t) \cdot \underbrace{p_\theta(X=x | T(X)=t)}_{\text{indp of } \theta.}$$

(for $t = T(x)$)

$$f_0 = h(x).$$

(see Keener textbook for a rigorous proof).

eg. Exponential family.

$$p_\theta(x) = \exp(\eta(\theta)^T T(x) - B(\theta)) \cdot h(x).$$

By factorization, T is sufficient

$\cdot N(\mu, \Sigma)$ in $\mathbb{R}^d$

$$p_{\mu, \Sigma}(x) = \frac{1}{(2\pi)^{d/2} \det(\Sigma)^{1/2}} \exp\left(-\frac{1}{2}(x-\mu)^T \Sigma^{-1}(x-\mu)\right)$$

$$\propto \exp\left(-\frac{1}{2} x^T \Sigma^{-1} x + \mu^T \Sigma^{-1} x\right)$$

$$\eta(\mu, \Sigma) = \begin{bmatrix} \Sigma^{-1} \mu \\ \text{vec}(\Sigma^{-1}) \end{bmatrix}$$

$$T(x) = \begin{bmatrix} x \\ \text{vec}(xx^T) \end{bmatrix} \in \mathbb{R}^{d+d^2}$$

(actually, we only need $d + \frac{d(d+1)}{2}$ dimensions).

- $X \sim \mathcal{N}(\mu, I_d)$. $T(x) = X \in \mathbb{R}^d$ gives an exp family
- $X \sim \mathcal{N}(0, \Sigma)$ $T(x) = \text{vec}(XX^T)$ gives an exp family
- $X \sim \text{Binom}(n, p)$ with known n .

$$p_p(x) = \binom{n}{x} \cdot p^x (1-p)^{n-x}$$

$$= \binom{n}{x} \cdot \exp\left(x \cdot \log \frac{p}{1-p} + n \log(1-p)\right).$$

$$T(x) = x. \quad \eta(p) = \log \frac{p}{1-p}.$$

$\eta(\theta)$: natural parametrization

Fact. $X_1, X_2, \dots, X_n \stackrel{\text{iid}}{\sim} p_\theta$ where $\{p_\theta: \theta \in \Theta\}$ is an exp family.

$$p_\theta(x) = \exp(\eta(\theta)^T T(x) - B(\theta)) \cdot h(x).$$

then $p_\theta(x_1, \dots, x_n) = \exp\left(\eta(\theta)^T \sum_{i=1}^n T(x_i) - n B(\theta)\right) \prod_{i=1}^n h(x_i).$

is also an exp family w/ $\sum_{i=1}^n T(x_i)$ sufficient.

Corollary. $X_1, \dots, X_n \stackrel{\text{iid}}{\sim} \mathcal{N}(\mu, \Sigma)$

then $(X_1, \dots, X_n)$ follows an exp family.
with $T(X) = \begin{bmatrix} \sum_{i=1}^n x_i \\ \sum_{i=1}^n \text{vec}(x_i x_i^T) \end{bmatrix}$

eg. $\text{Unif}([\theta-1, \theta+1])$ $\theta \in \mathbb{R}$.

$X_1, X_2, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} \text{Unif}([\theta-1, \theta+1])$.

$$T(X) = (X_{(1)}, X_{(n)})$$

where $X_{(1)}, X_{(2)}, \dots, X_{(n)}$

are sorted version of $X_1, \dots, X_n$

$$P_\theta(X_1, \dots, X_n) = 2^{-n} \cdot \mathbb{1}_{\{X_{(1)} \geq \theta-1, X_{(n)} \leq \theta+1\}}.$$

Use of sufficient stats: unbiased estimation.

Estimate $g(\theta)$ from $X \sim P_\theta$.

Def. δ is unbiased for $g(\theta)$ if

$$\mathbb{E}_\theta[\delta(X)] = g(\theta) \quad (\forall \theta \in \Theta)$$

Def. (UMVU) For any other δ' unbiased.

$$\text{we have } \text{var}_\theta(\delta(X)) \leq \text{var}_\theta(\delta'(X)) \quad (\forall \theta \in \Theta)$$

Among all the possible unbiased estimators, δ is the optimal.

Def (complete stats)

We call $T(X)$ is complete for $\{P_\theta: \theta \in \Theta\}$ if

$$\mathbb{E}_\theta[f(T(X))] = 0, \forall \theta \in \Theta \quad \text{implies} \quad f(T) = 0 \quad \text{a.s.}$$

A system of integral eq.
on the function f .

only have
trivial solutions.

Idea: no redundancy of info in T .
eg. If $X_1, X_2, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} \text{Unif}(\theta-1, \theta+1)$.
 $T(X_1, \dots, X_n) = (X_1, \dots, X_n)$.

$$f(T) = x_1 - x_2$$

But for $T(X_1, \dots, X_n) = (X_{(1)}, X_{(n)})$.
you cannot find such functions.

eg. $X_1, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} \text{Unif}([0, \theta])$.

$$T(X) = \max_{1 \leq i \leq n} X_i. \quad \mathbb{P}_\theta(T \leq t) = \left(\frac{t}{\theta}\right)^n \quad \text{for } t \in [0, \theta].$$

$$\mathbb{E}_\theta[f(T(X))] = \frac{n}{\theta^n} \int_0^\theta f(t) t^{n-1} dt = 0 \quad (\forall \theta).$$

$$\Rightarrow \int_0^\theta f(t) t^{n-1} dt = 0 \quad \forall \theta$$

So. $f(t) = 0$ almost everywhere.

Def. Exp family $p_{\theta}(x) = \exp(\eta(\theta)^T T(x) - B(\theta)) h(x)$
 is called full-rank if $\eta(\Theta)$ has an interior point.
 $\left(T_1(x), T_2(x), \dots, T_d(x) \right)$
 linearly indep.

Fact. Full-rank exp family $\Rightarrow T$ is sufficient & complete.

Proof. Let $\eta(\theta_0)$ be an interior point of $\eta(\Theta)$.

Let ν be density of $T(X)$ under θ_0 .

(Idea: local perturbations around θ_0 .)

$$0 = \int_{\mathbb{R}^d} f(t) \cdot \frac{p_{\theta}(T(X)=t)}{\nu(t)} \cdot \nu(t) dt$$

$$\Rightarrow 0 = \int_{\mathbb{R}^d} f(t) \cdot \exp((\eta(\theta) - \eta(\theta_0))^T t) \nu(t) dt \quad (\forall \theta \in \Theta).$$

$$\text{So. } \int_{\mathbb{R}^d} f(t) e^{\eta^T t} \nu(t) dt = 0$$

for $\eta \in B(0, r_0)$

multivariate
Laplace transform
of $f \cdot \nu$.

By uniqueness of Laplace transform.
 $f \cdot \nu = 0$ (a.e.).

Thm (Lehmann - Scheffé)

If T is sufficient & complete for $(\mathcal{P}_\theta: \theta \in \Theta)$

Suppose that $\exists$ unbiased estimator δ for $g(\theta)$.

then $\delta^* = E[\delta(X)|T]$ is UMVU.

Proof: δ^* is unbiased.

Consider another unbiased estimator δ' .

$$\tilde{\delta}(T) = E[\delta'(X)|T]$$

$$\forall \theta, \text{Var}_\theta(\tilde{\delta}) \leq \text{Var}(\delta').$$

On the other hand, by completeness.

$$E_\theta[\tilde{\delta}(T) - \delta^*(T)] = 0 \quad (\forall \theta \in \Theta).$$

$$\text{So } \tilde{\delta} = \delta^* \quad (\text{a.s.}).$$

e.g. $X_1, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} N(\theta, 1)$. $g(\theta) = \theta$.

$$\delta(X) = T(X) = \frac{1}{n} \sum_{i=1}^n X_i \quad \text{is UMVU.}$$

eg. $X_1, \dots, X_n \stackrel{\text{iid}}{\sim} \text{Ber}(\theta)$ $g(\theta) = \theta^k$

$$T(X) = \sum_{i=1}^n X_i.$$

Unbiased estimator

$$g(X) = X_1 \cdot X_2 \cdots X_k.$$

$$\text{UMVU: } E[g(X) | T(X)] = \frac{T(T-1) \cdots (T-k+1)}{n(n-1) \cdots (n-k+1)}.$$

eg. $X_1, \dots, X_n \stackrel{\text{iid}}{\sim} N(0, \sigma^2)$ $g(\sigma^2) = \sigma^2.$

$$T(X) = \sum_{i=1}^n X_i^2$$

$$\frac{T(X)}{n} \text{ is UMVU.}$$

Indeed, UMVU is not the "optimal" estimator

$$\delta_c(T) = cT.$$

$$\arg \min_c \text{MSE}(c) = \frac{1}{n+2}.$$

Additional remarks. (sufficiency, unbiased estimation)

Least useful

most useful

UMVU

Completeness

sufficiency

Rao-Blackwell.

debiasing

(approx unbiased).

Cramér-Rao.

Thm. (Cramér-Rao)

If $g(X)$ is unbiased estimator for $g(\theta) \in \mathbb{R}$.

Assuming $(p_\theta: \theta \in \Theta)$ has shared support $(\forall \theta)$.

$$\log p_\theta(x) \in C^2, \quad \mathbb{E} |\nabla \log p_\theta(x)|^2 < +\infty$$

then we have *twice as differentiable*

$$\text{Var}_\theta(g(X)) \geq \nabla g(\theta)^T I(\theta)^{-1} \nabla g(\theta)$$

$$\text{where } I(\theta) := \mathbb{E}_\theta [\nabla \log p_\theta(x) \cdot \nabla \log p_\theta(x)^T].$$

(usually not achieved by unbiased estimation, achievable (asymptotically) by MLE plug-in).

Proof. two simple facts.

$$l(\theta; X) = \log p_\theta(X).$$

$$\begin{aligned} \mathbb{E}_\theta [\nabla_\theta l(\theta; X)] &= \int \underbrace{(\nabla_\theta l(\theta; x) \cdot e^{l(\theta; x)})}_{\nabla_\theta e^{l(\theta; x)}} d\mu(x). \\ &= \nabla_\theta \left(\underbrace{\int e^{l(\theta; x)} d\mu(x)}_{=1} \right) = 0. \end{aligned}$$

• (Fisher's identity).

Note that $\mathbb{E}_{\theta} \left[\frac{\partial}{\partial \theta_j} \ell(\theta; x) \right] = 0.$

take more derivative.

$$\begin{aligned} 0 &= \frac{\partial}{\partial \theta_i} \mathbb{E}_{\theta} \left[\frac{\partial}{\partial \theta_j} \ell(\theta; x) \right] \\ &= \frac{\partial}{\partial \theta_i} \int \frac{\partial}{\partial \theta_j} \ell(\theta; x) \cdot e^{\ell(\theta; x)} d\mu(x) \\ &= \int \left(\frac{\partial^2 \ell(\theta; x)}{\partial \theta_j \partial \theta_i} + \frac{\partial \ell(\theta; x)}{\partial \theta_i} \cdot \frac{\partial \ell(\theta; x)}{\partial \theta_j} \right) e^{\ell(\theta; x)} d\mu(x) \end{aligned}$$

So we have

$$\mathbb{E}_{\theta} \left[\nabla \ell(\theta; x) \nabla \ell(\theta; x)^T \right] = - \mathbb{E}_{\theta} \left[\nabla^2 \ell(\theta; x) \right].$$

(holding true only for well specified model at true θ .)

Proof of CRLB.

$$\begin{aligned} \nabla_{\theta} g(\theta) &\stackrel{(\text{unbias})}{=} \nabla_{\theta} \left(\int \delta(x) \cdot e^{\ell(\theta; x)} d\mu(x) \right) \\ &= \int \delta(x) \cdot \nabla_{\theta} \ell(\theta; x) \cdot e^{\ell(\theta; x)} d\mu(x). \end{aligned}$$

How to insert $g(\theta)$?

$$\int \nabla_{\theta} \ell(\theta; x) e^{\ell(\theta; x)} d\mu(x) = 0.$$

$g(\theta)$ indep of x .

$$\begin{aligned}\nabla g(\theta) &= \int (\delta(x) - g(\theta)) \cdot \nabla_{\theta} \ell(\theta; x) \cdot e^{\ell(\theta; x)} d\mu(x) \\ &= \mathbb{E}_{\theta} [(\delta(x) - g(\theta)) \cdot \nabla_{\theta} \ell(\theta; x)],\end{aligned}$$

in 1D. Cauchy-Schwarz gives the result.

in general, consider quadratic form

$$u \in \mathbb{R}^d \mapsto \mathbb{E}_{\theta} \left[\left(u^T \nabla_{\theta} \ell(\theta; x) - (\delta(x) - g(\theta)) \right)^2 \right] \geq 0.$$

Expanding the quadratic form.

$$0 \leq u^T I(\theta) u - 2 u^T \mathbb{E}_{\theta} [(\delta(x) - g(\theta)) \cdot \nabla_{\theta} \ell(\theta; x)] + \text{var}_{\theta}(\delta(x)) = \nabla g(\theta).$$

This implies $\text{var}_{\theta}(\delta(x)) \geq \nabla g(\theta)^T I(\theta)^{-1} \nabla g(\theta)$.

Corollary. $g: \mathbb{R}^d \rightarrow \mathbb{R}^k$. δ unbiased.

$$\mathbb{E} \left[(\delta(x) - g(\theta))^{\otimes 2} \right] \succeq_{(PSD)} (\nabla g(\theta))^T I(\theta)^{-1} (\nabla g(\theta)).$$

In particular, $g(\theta) = \theta$

$$\mathbb{E} [\|\delta(x) - \theta\|_2^2] \geq \text{Tr}(I(\theta)^{-1}).$$

eg. $X_1, \dots, X_n \stackrel{\text{iid}}{\sim} p_\theta$

$$l_n(\theta; X_1, \dots, X_n) = \sum_{i=1}^n l(\theta; X_i).$$

$$I_n(\theta) = n \cdot I(\theta).$$

and CRLB
$$\frac{(\nabla g(\theta))^T I(\theta)^{-1} \nabla g(\theta)}{n}$$

 (asymptotically achieved by MLE, $n \rightarrow \infty$)

When $p_\theta(x)$ is singular / has discontinuities

CRLB is false.

eg. $X_1, X_2, \dots, X_n \stackrel{\text{iid}}{\sim} \text{Unif}([0, \theta]).$

UMVU.

$$\delta(X) = \max_i X_i \cdot \frac{n+1}{n}.$$

$$\text{Var}_\theta(\delta(X)) = \frac{\theta^2}{n(n+2)} \quad \text{faster rate.}$$

Bayesian CRLB.

Idea. key step in CRLB

$$\nabla g(\theta) = \int (\delta(X) - g(\theta)) \nabla l(\theta; X) \cdot e^{l(\theta; X)} d\mu(X).$$

relying on unbiased of δ .

"look like" integration-by-parts formula.

use a prior distribution to integrate.

$$r_{\pi}(\delta) := \int_{\Theta} \mathbb{E}_{\theta}[(\delta(x) - g(\theta))^2] \pi(d\theta).$$

Thm (van Trees inequality). for any estimator

$$r_{\pi}(\delta) \geq \left(\int_{\Theta} \nabla g(\theta) \pi(d\theta) \right)^T \left(\int_{\Theta} I(\theta) \pi(d\theta) + J(\pi) \right)^{-1} \cdot \left(\int_{\Theta} \nabla g(\theta) \pi(d\theta) \right).$$

where $J(\pi) := \int_{\Theta} \nabla_{\theta} \log \pi(\theta) \nabla_{\theta} \log \pi(\theta)^T \pi(d\theta)$

Information theory's Fisher info".

Compared to CRLB:

- Pay a $J(\pi)$ term
- No need for unbiasedness.

Proof: see next lecture.

eg. $v_0(x) = \cos^2\left(\frac{\pi x}{2}\right) \mathbb{1}_{|x| \leq 1}$

supported on $[-1, 1]$.

$$J(v_0) = \pi^2.$$