

Midterm: 10/17, 5:10pm - 7:40pm
(2 hrs and 30 min)
in the classroom.

- 4 pages (double-sided) cheat sheet
 - no electronics.
 - Bring your student ID.
 - Covers first 6 lectures
 - Format. (4 questions in total)
 1. True or false
 2. 3. 4. Calculations/proofs/descriptions.
(show your reasoning).
(multiple parts).
- sufficient
startbooks

From last lecture: consistency of M-estimators (incl. MLE).

$$\hat{\theta}_n = \underset{\theta \in \Theta}{\operatorname{argmin}} F_n(\theta) \quad \text{where } F_n(\theta) := \frac{1}{n} \sum_{i=1}^n f(\theta; x_i)$$

$$F(\theta) := \mathbb{E}[f(\theta; X)], \quad \theta^* := \underset{\theta \in \Theta}{\operatorname{argmin}} F(\theta).$$

(for MLE, $f(\theta; x) = -\log p_\theta(x)$).

$$\hat{\theta}_n \xrightarrow{P} \theta^*.$$

Holds true under the following conditions.

(i) (ULLN). $\sup_{\theta \in \Theta} |F_n(\theta) - F(\theta)| \xrightarrow{P} 0.$

(ii) $\forall \varepsilon > 0, \exists \delta > 0$
st. $F(\theta) \geq F(\theta^*) + \varepsilon$
when $\|\theta - \theta^*\| \geq \delta.$

(Equivalently. $\inf_{\theta: \|\theta - \theta^*\| \geq \delta} F(\theta) > F(\theta^*) \quad \forall \delta > 0$)

(iii) Holds true when F is a continuous function, Θ is compact.
and θ^* is unique minimizer

(for finite-dim models, compactness $\Leftrightarrow$ bounded and closed)

Proof: F is cts function on the domain

$$\Theta \setminus \{\theta : \|\theta - \theta^*\| < \delta\}.$$

So $\inf_{\substack{\|\theta - \theta^*\| \geq \delta \\ \theta \in \Theta}} F(\theta)$ is attained by some θ_0

since θ^* is unique m.m., $F(\theta_0) > F(\theta^*)$.

When specialized to MLE:

$$F(\theta) = -\mathbb{E}_{\theta^*}[\log p_{\theta}(X)]$$

$$= \underbrace{\mathbb{E}_{\theta^*}[\log p_{\theta^*}(X)]}_{\text{const, indep of } \theta} + \underbrace{\mathbb{E}_{\theta^*}\left[\log \frac{p_{\theta^*}(X)}{p_{\theta}(X)}\right]}_{\text{Kullback-Leibler divergence.}}$$

$$D_{KL}(P \parallel Q) = \mathbb{E}_P\left[\log \frac{p(X)}{q(X)}\right]$$

where p, q are densities of P, Q .

- $D_{KL}(P \parallel Q) \geq 0 \quad (\forall P, Q)$

$$D_{KL}(P \parallel Q) \geq \frac{1}{2} \left(\int |p(x) - q(x)| dx \right)^2$$

(Pinsker's inequality).

$$D_{KL}(P \parallel Q) = 0 \quad \text{if and only if} \quad P = Q.$$

For MLE

$$F(\theta) = \text{constant} + D_{KL}(\mathbb{P}_{\theta^*} \parallel \mathbb{P}_{\theta})$$

- minimized at $\theta = \theta^*$

- When the model is "identifiable", θ^* uniquely minimizes F .
(If $\theta_1 \neq \theta_2$ then $\mathbb{P}_{\theta_1} \neq \mathbb{P}_{\theta_2}$)

On identifiability:

eg. $\mathbb{P}_{\theta} = \frac{1}{2} \mathcal{N}(\theta, 1) + \frac{1}{2} \mathcal{N}(-\theta, 1)$.

$\theta \in \mathbb{R}$, then it's not identifiable

$$\mathbb{P}_{\theta} = \mathbb{P}_{-\theta}$$

$\Theta = \{\theta : \theta \geq 0\}$, the model is identifiable.

Still need to show (i)

Thm (Wald) Suppose Θ is compact
(closed and bounded in $\mathbb{R}^d$)

Assuming:

$$\cdot \mathbb{E} \left[\sup_{\theta \in \Theta} |f(\theta; X)| \right] < +\infty \quad (\text{from DCT})$$

$\cdot f(\theta; X)$ is a cts function of θ , $\forall X$.

$$\text{then } \sup_{\theta \in \Theta} |F_n(\theta) - F(\theta)| \xrightarrow{P} 0.$$

Remark: . Holds true under very general setting.

Θ doesn't even need to be finite-dim.

• Intuition:

If Θ is finite,

$$\mathbb{P} \left(\sup_{\theta \in \Theta} |F_n(\theta) - F(\theta)| > \varepsilon \right)$$

$$\leq \sum_{\theta \in \Theta} \mathbb{P} (|F_n(\theta) - F(\theta)| > \varepsilon) \rightarrow 0$$

Compactness + continuity allows us
to use discrete approximations to Θ .

Implications for MLE:

identifiability + obs log-likelihood + $\mathbb{E}[\sup_{\theta \in \Theta} l(\theta)] < \infty$

$$\Rightarrow \hat{\theta}_n \xrightarrow{P} \theta^*$$

furthermore, (from last lecture)
under additional smoothness assumptions & $I(\theta^*) > 0$.

then
$$\sqrt{n}(\hat{\theta}_n - \theta^*) \xrightarrow{d} \mathcal{N}(0, I(\theta^*))$$

(since we can show MLE consistency,
we don't need to restrict in a local neighborhood
of θ^*).

Jackknife.

Motivation: in many cases, the estimator
satisfies asymptotic expansion

$$\hat{\theta}_n - \theta^* = \frac{1}{n} \sum_{i=1}^n z_i + \frac{a}{n} + \frac{b}{n^2} + O_p\left(\frac{1}{n^3}\right)$$

$$(z_i \stackrel{iid}{\sim}, E[z] = 0, E[z]^2 < +\infty)$$

$$\left| \frac{1}{n} \sum_{i=1}^n z_i \right| = O\left(\frac{1}{\sqrt{n}}\right)$$

Jackknife bias estimator

Let $\hat{\theta}_{(-i)}$ be the estimator computed from $(X_j)_{j \neq i}$. ($i=1, 2, \dots, n$)

$$\hat{\theta}_{(-i)} - \theta^* = \frac{1}{n-1} \sum_{j \neq i} z_j + \frac{a}{n-1} + \frac{b}{(n-1)^2} + O_p\left(\frac{1}{n^3}\right)$$

"leave-one-out" estimators.

$$\begin{aligned} \hat{b}_{jack} &= \frac{n-1}{n} \sum_{i=1}^n \hat{\theta}_{(-i)} - (n-1) \hat{\theta}_n \\ &= \left(\cancel{\frac{n-1}{n} \sum_{i=1}^n z_i} + \frac{(n-1)}{n} \cdot n \cdot \frac{a}{n-1} + O_p\left(\frac{1}{n^2}\right) \right) \text{ "ignore"} \\ &\quad - \left(\cancel{\frac{n-1}{n} \sum_{i=1}^n z_i} + \frac{n-1}{n} \cdot a + O_p\left(\frac{1}{n^2}\right) \right) \end{aligned}$$

$\hat{b}_{jack}$ as estimator for $\frac{a}{n}$ difference

Jackknife debiasing.

$$\hat{\theta}_{\text{debias}} = \hat{\theta}_n - \hat{b}_{\text{jack}}.$$

- useful in some high-dim applications.
 - expansion holds true only w/ $n \rightarrow +\infty$
for finite n , $\hat{\theta}_{\text{debias}}$ may still be biased.
-

Computational algorithms.

- Newton's method.

$$\min_{\theta \in \mathbb{R}^d} \{F(\theta)\}.$$

Start from $\theta^{(0)}$, iterate

$$\theta^{(t+1)} = \theta^{(t)} - \nabla^2 F(\theta^{(t)})^{-1} \nabla F(\theta^{(t)})$$

($t = 0, 1, 2, \dots$).

- Local super-exponential convergence.

One-step Newton estimator:

Start from $\tilde{\theta}_n$

$$\hat{\Theta}_n := \tilde{\Theta}_n - \nabla^2 F_n(\tilde{\Theta}_n)^{-1} \cdot \nabla F_n(\tilde{\Theta}_n).$$

Fact. If $\|\tilde{\Theta}_n - \theta^*\| = O_p\left(\frac{1}{\sqrt{n}}\right)$

(this holds true when $\sqrt{n}(\tilde{\Theta}_n - \theta^*) \xrightarrow{d} \mathcal{N}(0, \Sigma_0)$
for some potentially large Σ_0).

then $\sqrt{n}(\hat{\Theta}_n - \theta^*) \xrightarrow{d} \mathcal{N}(0, H_*^{-1} \Sigma_* H_*^{-1})$

where $H_* = \nabla^2 F(\theta^*)$

$$\Sigma_* = \text{cov}(\nabla f(\theta^*, X))$$

(the same limiting distribution as M-estimator)

Proof... $\tilde{\Theta}_n - \nabla^2 F_n(\tilde{\Theta}_n)^{-1} \nabla F_n(\tilde{\Theta}_n)$

$$\approx \underbrace{\tilde{\Theta}_n - \nabla^2 F_n(\theta^*)^{-1} \cdot \nabla F_n(\theta^*)}_{\text{cancellation}} + \underbrace{\nabla^2 F_n(\theta^*)^{-1} \cdot \nabla^2 F_n(\theta^*) \cdot (\tilde{\Theta}_n - \theta^*)}_{\text{cancellation}}$$

$$\approx \underbrace{\theta^* - \nabla^2 F(\theta^*)^{-1} \cdot \nabla F_n(\theta^*)}_{\text{cancellation}}$$

$$\sqrt{n}(\dots) \xrightarrow{d} \mathcal{N}(0, H_*^{-1} \Sigma_* H_*^{-1}).$$

• SGD.

$$\theta_{t+1} = \theta_t - \alpha_{t+1} \nabla f(\theta_t; X_{t+1})$$

(for $t=0, 1, 2, \dots, n-1$)

$(\alpha_t)_{t \geq 1}$ step sequence (tuning parameters)

(Polyak - Juditsky - Ruppert)

$$\hat{\theta}_n := \frac{1}{n} \sum_{t=1}^n \theta_t.$$

Fact. under certain regularity conditions (incl convexity)
with suitable stepsize sequence,

$$\sqrt{n}(\hat{\theta}_n - \theta^*) \xrightarrow{d} \mathcal{N}(0, H_{\theta^*}^{-1} \Sigma_{\theta^*} H_{\theta^*}^{-1})$$

(same as M-estimator)

Sufficient statistics.

Def. $(P_{\theta} : \theta \in \Theta)$ is a class of distributions

We call $T(X)$ sufficient if
 $\forall \theta$ and $\forall t$, conditional distribution of X

under P_θ conditioned on $T(X)=t$

's indep of θ .

eg. $X_1, X_2, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} \text{Ber}(p)$.

$$T(X_1, X_2, \dots, X_n) = X_1 + X_2 + \dots + X_n \sim \text{Binom}(n, p)$$

Conditionally on $T(X_1, \dots, X_n)=t$

$(X_1, \dots, X_n)$ are indicators of random subset
with cardinality t , indep of p .

Thm (factorization)

$$T \text{ is sufficient } \Leftrightarrow \exists g_\theta, h$$
$$\text{st. } P_\theta(x) = g_\theta(T(x)) \cdot h(x)$$
$$(\forall \theta, x)$$

eg. $X_1, X_2, \dots, X_n \stackrel{\text{i.i.d.}}{\sim} N(\mu, 1)$

$$P_\mu(x_1, \dots, x_n) = \frac{1}{(2\pi)^{n/2}} \exp\left(-\frac{1}{2} \sum_{i=1}^n (\mu - x_i)^2\right)$$

$$= \frac{1}{(2\pi)^{n/2}} \cdot \exp\left(-\mu \cdot \left(\sum_{i=1}^n x_i\right) - \frac{n\mu^2}{2}\right) \cdot \exp\left(-\frac{1}{2} \sum_{i=1}^n x_i^2\right)$$

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

Proof. " $\Leftarrow$ "

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

indp of θ .

" $\Rightarrow$ " Construction: $g_{\theta}(t) = P_{\theta}(T(X)=t)$
(density of $T(X)$ when $X \sim P_{\theta}$)

$h(x) = P_{\theta_0}(X=x|T(X)=t) (= P_{\theta}(X=x|T(X)=t))$
(conditional density of X conditioned on $T(X)=t$
under model $X \sim P_{\theta_0}$).

Fact. For any estimator $\hat{\theta}_n$ of θ .

$\exists$ another estimator $\tilde{\theta}_n$ which depends on X
only through $T(X)$

such that $\hat{\theta}_n \stackrel{d}{=} \tilde{\theta}_n$ under P_{θ} ($\forall \theta$).

Proof: $X|T(X)$ is indep of θ .

So we can sample from this conditional distribution.

- Sample $\tilde{X} \sim p(x|T(x)=t)$
- Compute $\hat{\theta}_n$ on $\tilde{X}$, and let $\tilde{\theta}_n$ be the result.

Improving an estimator using sufficient stats.

• Rao-Blackwellization.

Given estimator $\hat{\theta}_n$ and sufficient stats $T(X)$

$$\hat{\theta}_{RB} := \mathbb{E}[\hat{\theta}_n | T(X)].$$

$$\bullet \quad \mathbb{E}_{\theta}[\hat{\theta}_{RB}] = \mathbb{E}_{\theta}[\hat{\theta}_n]$$

(If $\hat{\theta}_n$ is unbiased, so is $\hat{\theta}_{RB}$).

$$\bullet \quad \text{Var}_{\theta}(\hat{\theta}_{RB}) \leq \text{Var}_{\theta}(\hat{\theta}_n).$$

Exponential family.

Suppose $X_1, X_2, \dots, X_n \stackrel{\text{iid}}{\sim} P_\theta$

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

For joint distribution

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

By factorization thm.

$\sum_{i=1}^n T(x_i)$ is sufficient for θ .

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

$$P_{\mu, \sigma^2}(x) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{x^2}{2\sigma^2} - \frac{\mu^2}{2\sigma^2} + \frac{x\mu}{\sigma^2}\right).$$

$$= \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{\mu^2}{2\sigma^2}\right) \cdot \exp\left(\begin{bmatrix} \mu/\sigma^2 \\ -\frac{1}{2\sigma^2} \end{bmatrix}^T \begin{bmatrix} x \\ x^2 \end{bmatrix}\right).$$

$$T(X_1, X_2, \dots, X_n) = \begin{bmatrix} \sum_{i=1}^n X_i \\ \sum_{i=1}^n X_i^2 \end{bmatrix} \text{ is sufficient}$$