

Model selection

1

- So far, assumed we were working with the true model family $\mathcal{P}$ $\rightarrow$ estimation problem just one of fixing parameter values.

What happens if we use the wrong model!

Return to Information inequality

Assume the model is g , and we use model f .

(Example g is t-distribution and f is normal)

Measure of model distance KL

$$KL(g, f) = E_g \log \frac{g(x)}{f(x)} = \int \left(\log \frac{g(x)}{f(x)} \right) g(x) dx \geq 0$$

(shown before)

With equality if $g = f$

$$\Rightarrow E_g \log g(x) \geq E_g \log f(x)$$

So expected log likelihood is greater under true model than wrong model (wrt truth)

What happens if we proceed with MLE using f ?

(2)

$$L_f(\theta | X^n) = \prod_{i=1}^n f_\theta(x_i) \Rightarrow \hat{\theta} = \underset{\theta}{\operatorname{argmax}} \frac{1}{n} \sum_{i=1}^n \log f_\theta(x_i)$$

$$\xrightarrow{\text{WLLN}} E_g[\log f_\theta(X)]$$

note g - since that's the data generating process

$$\text{So } \max_{\theta} L_f(\theta | X^n) \stackrel{\text{large } n}{=} \min_{\theta} (-E_g[\log f_\theta(X)])$$

$$= \min_{\theta} \underbrace{E_g \log g(X) - E_g[\log f_\theta(X)]}_{\text{unknown (constant)}}$$

$$= \min_{\theta} KL(g, f)$$

$\Rightarrow$ so MLE equivalent to finding closest f to g among $\{f_\theta : \theta \in \Theta\}$.

$$\text{As } n \rightarrow \infty \quad \hat{\theta} = \underset{\theta}{\operatorname{argmax}} L_f(\theta | X^n) \Rightarrow \text{converges to } \underset{\theta}{\operatorname{argmax}} E_g[\log f_\theta(X)] = \theta^*$$

Where f_{θ^*} closest model to the model g .

Note, θ^* may have nothing to do with g itself.

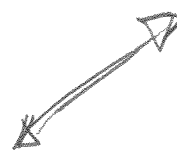
When parameter θ , and estimate $\hat{\theta} \rightarrow \theta^*$ where θ^* is meaningful i.e. θ^* corresponds to some population parameters, like mean & variance, under g as well, the modeling is called robust.

(Example from book $g \sim \text{Gamma}(4, 1)$ mean $4 \cdot 1 = 4$
variance $4 \cdot 1^2 = 4$)

$$f \sim N(0, \sigma^2)$$

$$\text{turns out } \hat{\theta} \rightarrow 4$$

$$\hat{\sigma}^2 \rightarrow 4$$



(but) no free lunch $\Rightarrow$ tends to be the case that $\hat{\theta}_f$ less efficient than $\hat{\theta}_g$.

Consistency result here

$$\theta_0 = \operatorname{argmax}_{\theta} E_g \log f(X)$$

$$\text{Under reg. conditions } \hat{\theta} = \operatorname{argmax}_{\theta} \bar{\ell}_f(\theta | X^n)$$

$$= \operatorname{argmax}_{\theta} \frac{1}{n} \sum_{i=1}^n \log f_{\theta}(x_i) \rightarrow^P \theta_0$$

(Similar proof as before)

What about other asymptotic results?

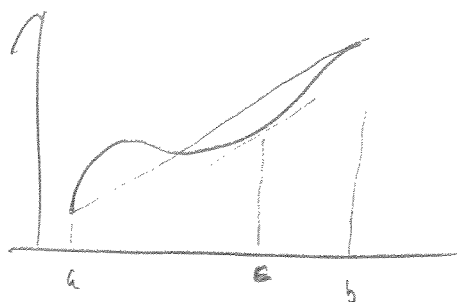
4

$$S_f(\theta) = S_{\text{score}} = \frac{\partial}{\partial \theta} \log L_f(\theta | X^n) = \frac{\partial}{\partial \theta} \sum_{i=1}^n \log f_{\theta}(x_i)$$

$$\hat{\theta}_f = \text{argmax}_{\theta} S_f(\theta)$$

$$\text{Expand score around } \hat{\theta} \Rightarrow S_f(\theta) \approx S_f(\hat{\theta}) + \frac{\partial^2}{\partial \theta \partial \theta'} \log L_f(\theta | X^n) \Big|_{\hat{\theta}} (\theta - \hat{\theta})$$

Mean-value Theorem



$$\frac{f(b) - f(a)}{b - a} = f'(c)$$

some $c \in [a, b]$

$\Rightarrow$ Here, take θ^* somewhere between θ and $\hat{\theta}$

$$\Rightarrow S_f(\theta) = S_f(\hat{\theta}) + \frac{\partial^2}{\partial \theta \partial \theta'} \log L_f(\theta | X^n) \Big|_{\theta^*} (\theta - \hat{\theta})$$

$$= \frac{\partial^2}{\partial \theta \partial \theta'} \log L_f(\theta | X^n) \Big|_{\theta^*} (\theta - \hat{\theta})$$

$$\text{Now } S_f(\theta) = \frac{\partial}{\partial \theta} \sum_{i=1}^n \log f_{\theta}(x_i) \quad \text{where } E(Y_i) = E_g \left[\frac{\partial}{\partial \theta} \log f_{\theta}(X) \right]$$

$$= \frac{\partial}{\partial \theta} E_g [\log f_{\theta}(X)] = \frac{\partial}{\partial \theta} \int (\log f_{\theta}(x)) g(x) dx$$

for nice $g \circ f$

and since $\theta_0 = \text{maximizer of } E_g(\log f) \Rightarrow E(Y_i) = 0 \text{ @ } \theta = \theta_0$

Moreover $\text{Var}(Y_i) = E \left[\left(\frac{\partial \log f_{\theta}}{\partial \theta} \right) \left(\frac{\partial \log f_{\theta}}{\partial \theta} \right)' \right] \Big|_{\theta=\theta_0} = J(\theta_0)$

and by CLT $\frac{1}{\sqrt{n}} \sum_{i=1}^n Y_i \rightarrow^d N(0, J)$

Now return to score

$$S_f(\theta) = \frac{\partial^2}{\partial \theta \partial \theta'} \log L_f(\theta | X^n) \Big|_{\theta^*} (\theta - \hat{\theta})$$

$$n \cdot \left(\frac{1}{n} \sum_{i=1}^n \frac{\partial^2}{\partial \theta \partial \theta'} \log f_{\theta}(x_i) \Big|_{\theta^*} \right)$$

$$\xrightarrow{P} E \left(\frac{\partial^2}{\partial \theta \partial \theta'} \log f_{\theta}(X) \right) \Big|_{\theta_0} = -F I(\theta_0)$$

Since $\hat{\theta} \rightarrow \theta_0$

and θ^* between $\hat{\theta}$ and θ_0 .

$$\Rightarrow \frac{1}{\sqrt{n}} S_f(\hat{\theta}) = -\frac{1}{n} \frac{\partial^2}{\partial \theta \partial \theta'} \log L_f(\theta | X^n) \Big|_{\theta^*} \cdot \sqrt{n} (\hat{\theta} - \theta_0)$$

$$\xrightarrow{d} N(0, J(\theta_0)) \quad \xrightarrow{P} F I(\theta_0)$$

$$\Rightarrow \sqrt{n} (\hat{\theta} - \theta_0) \rightarrow^d N(0, F I^{-1}(\theta_0) J(\theta_0) F I^{-1}(\theta_0))$$

$= F I^{-1}(\theta_0) \quad \text{if } f_{\theta_0} = q$

So, asymptotic variance is $\underline{F^{-1}(\theta_0) J(\theta_0) F^{-1}(\theta_0)}$

6

(note, function of true unknown θ_0
and unknown g')

is of course, not so practical to use
 $\Rightarrow$ but can examine for model classes
to get some idea about (loss of)
efficiency

Breakdown of ML Principle

General idea was that maximizing likelihood $\Rightarrow$
a 'most likely' to generate data observed

"Bayesian
view"

BUT, now assume we explore a class of
models $\{f \in \mathcal{F}\}$ with parameter dimension $\{p_f, f \in \mathcal{F}\}$

Simple case: nested models $f_1 \subset f_2 \subset f_3 \dots \in f_n$

e.g.

regression models w.

$\beta_1, (\beta_1, \beta_2), (\beta_1, \beta_2, \beta_3) \dots$