

Model Selection

①

① Via testing

② Via generalization criteria, e.g. C_p , AIC or cross-validation

Testing

LRT (likelihood ratio test) see ch.9

Logic behind this $\rightarrow$ assumption of a correct null model and look for evidence against it.

Example $\theta = (\theta_1, \theta_2)$, $\theta_1 = \text{parameter(s) of interest}$, $\dim(\theta_1) = p$.

• Arrive at test via e.g. profile likelihood

where we replace θ_2 by the MLE for every value of $\theta_1 \rightarrow \hat{\theta}_2(\theta_1)$

$\Rightarrow$ Then focus on estimation problem of θ_1

$$\Rightarrow L(\theta_1) = \max_{\theta_2} L(\theta_1, \theta_2) = L(\theta_1, \hat{\theta}_2(\theta_1))$$

$\Rightarrow$ Now compare $\max_{\theta_1} L(\theta_1) = L(\hat{\theta}_1)$ to $L(\theta_{1,0})$ null $H_0: \theta_1 = \theta_{1,0}$

$$\Rightarrow \text{Will's stat } W = 2 \log \frac{L(\hat{\theta}_1)}{L(\theta_{1,0})} \sim \chi^2_{p,1} \text{ under null}$$

• (can also think of this as comparing $L(\theta_1, \theta_2)$ under max with restrictions on θ_1 or no restrictions

$$\Rightarrow W = 2 \log \frac{\max_{\theta_1, \theta_2} L(\theta_1, \theta_2)}{\max_{\theta_1 = \theta_{1,0}, \theta_2} L(\theta_1, \theta_2)} \sim \chi^2_{p,1}$$

General LRT result

$$W = 2 \log \frac{\max L(\theta)}{\max_{H_0} L(\theta)}, \quad \text{null } H_0 \text{ puts restrictions on } \theta \text{ (e.g. } \theta_1 = \theta_{10})$$

$$W \sim \chi^2_{\Delta_p} \text{ under } H_0, \text{ where } \Delta_p = \# \text{ restricted } \theta \text{ under } H_0.$$

- Result follows from Quad approx of $\log L$ and Asympt. Normality of $\hat{\theta}$, and Conditional Normality of $\hat{\theta}(\hat{\theta}(H_0))$
- (see Chapter 9)

Let's focus on AIC derivation

Some previous results we need;

$$L_f(\theta | x^n) = \prod_{i=1}^n f_\theta(x_i), \quad l_f(\theta) = \log L_f(\theta | x^n) = \sum_{i=1}^n \log f_\theta(x_i)$$

(log) likelihood under model f_θ (which may be incorrect)

$$\Rightarrow \hat{\theta}_f = \arg \max l_f(\theta), \text{ maximum value } l_f(\hat{\theta}_f)$$

If we now consider models $f_j, j=1, \dots, J$

where $\{f_j = f_{\theta_j}, \theta_j \in \Theta_j\}$, $p_j = \dim(\Theta_j)$

$\Rightarrow$ we cannot use $l_{f_j}(\hat{\theta}_j)$ to choose between f_j 's.

(3)

Reason why $l_f(\hat{\theta}_f)$ doesn't work $\rightarrow$ always
max for model f with most free parameters $\dim(\theta_f)$.

$\Rightarrow$ also $l_f(\hat{\theta}_f)$ only says something about the
fit on this data set x^n , but we want to
know if f works on future data as well
-generalization.

$\Rightarrow$ So, a more reasonable mode of comparison
is to choose the model f that
maximizes $E_g[\log L_f(\theta_f | X)]$
 $\uparrow$
new data x^{dg}

But, this doesn't work since θ_f is unknown

so use

$$E_g[\log L_f(\hat{\theta}_f(x^n) | X)] \equiv Q_f$$

$\uparrow$
expectation over x^n = training data used
to estimate θ_f

AND

x = test data, used to evaluate
how well f_{θ_f} works on future data

So - how to get a good estimate of

$$Q_f = E[\log L_f(\hat{\theta}_f(X^n) | X)] \quad ?$$

E over X and X^n

Observable: $\frac{1}{n} \ell_f(\hat{\theta}_f(X^n)) = \frac{1}{n} \sum_{i=1}^n \log f_{\hat{\theta}_f}(x_i) \quad \textcircled{*}$

is not a good estimate since $X = X^n$ used for both training & testing. $\textcircled{*}$ underestimates Q_f .

How big the bias is in $\textcircled{*}$ is a function of $\dim(\Theta_f)$.

Properties we will need later

We have $\hat{\theta}_f = \underset{\Theta_f}{\operatorname{argmax}} \ell_f(\theta_f | X^n)$

We have $\frac{1}{n} \sum_{i=1}^n \log f_{\hat{\theta}_f}(x_i) \rightarrow E_g[\log f_{\theta_f}(X)]$ by WLLN

and max of LHS $\rightarrow$ max of RHS

Say θ_0 maximizes $E_g[\log f_{\theta_f}(X)]$

then $\hat{\theta}_f \rightarrow \theta_0$

θ_0 may not be a 'true parameter' in the sense of true g .

Now $\hat{\theta}_f = \underset{\theta_f}{\operatorname{argmax}} l_f(\theta_f | X^n)$ or equivalently

(5)

The zero-crossing of the score $S_f(\theta) = \sum_{i=1}^n \frac{\partial}{\partial \theta} \log f_{\theta_f}(x_i)$

$$= \sum_{i=1}^n \frac{\partial}{\partial \theta} \log f_{\theta_f}(x_i)$$

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

(from mean-value theorem)

approx near $\hat{\theta}_f$

//
0
by
definition

2nd derivative of $\log L_f$
evaluated at θ^*

where

$$\|\theta - \theta^*\| \leq \|\theta_f - \hat{\theta}_f\| \quad (\text{somewhere in between } \theta \text{ and } \hat{\theta})$$

(I drop notation θ_f below and just use θ so it looks less messy)

OK, so

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

$$= \sum_{i=1}^n \underbrace{\frac{\partial}{\partial \theta} \log f_{\theta}(x_i)}_{y_i}$$

$$, y_i = \frac{\partial}{\partial \theta} \log f_{\theta}(x_i) \quad \text{random variable}$$

(6)

So the score $S_f(\theta)$ is a sum of iid random variables $Y_i \rightarrow$ CLT applies. Need to figure out what mean & variance is of the asymptotic distribution of $S_f(\theta)$.

- $E(Y_i) = E_g \left[\frac{\partial}{\partial \theta} \log f_{\theta}(X) \right] = \frac{\partial}{\partial \theta} E_g [\log f_{\theta}(X)]$

- $= 0$ at θ_0 since θ_0 is the maximizer of $E_g [\log f_{\theta}(X)]$

- At θ_0 , $\text{Var}(Y_i) = E(Y_i^2) = E_g \left[\left(\frac{\partial}{\partial \theta} \log f_{\theta}(X) \right) \left(\frac{\partial}{\partial \theta} \log f_{\theta}(X) \right)' \right]_{\theta_0}$
 $= J(\theta_0)$

- By CLT, $\frac{1}{\sqrt{n}} S_f(\theta_0) \rightarrow^d N(0, J(\theta_0))$

We have

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

$\rightarrow N(0, J(\theta_0))$

Went distribution of $\hat{\theta}$

Focus on $\frac{\partial^2}{\partial \theta \partial \theta'} \log L_f(\theta | X^n) \Big|_{\theta^*} \stackrel{a.s.}{\rightarrow} \textcircled{7}$

(7)

Now, since $\|\theta - \hat{\theta}\| \leq \|\theta - \tilde{\theta}\|$ and $\hat{\theta} \rightarrow \theta_0$

at θ_0 as $n \rightarrow \infty$ $\textcircled{7} \approx \frac{\partial^2}{\partial \theta \partial \theta'} \log L_f(\theta | X^n) \Big|_{\theta_0}$

Also $\frac{1}{n} \frac{\partial^2}{\partial \theta \partial \theta'} \log L_f(\theta | X^n) \Big|_{\theta_0}$

$= \frac{1}{n} \sum_{i=1}^n \frac{\partial^2}{\partial \theta \partial \theta'} \log f_{\theta}(x_i) \Big|_{\theta_0}$

$\rightarrow E \left[\frac{\partial^2}{\partial \theta \partial \theta'} \log f_{\theta}(X) \Big|_{\theta_0} \right]$ by WLLN

$= -FI(\theta_0)$ by definition

$\frac{1}{\sqrt{n}} S_f(\theta_0) \approx \left[\frac{1}{n} \sum_{i=1}^n \frac{\partial^2}{\partial \theta \partial \theta'} \log f_{\theta}(x_i) \Big|_{\theta_0} \right] \sqrt{n} (\theta_0 - \hat{\theta})$

$\rightarrow N(0, J(\theta_0))$

$\rightarrow -FI(\theta_0)$

By Slutsky's

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

Example

8

Assume $f_\lambda(x) = \lambda e^{-\lambda x}$, $\lambda > 0$ (exponential)
 $x \geq 0$

$$\log f_\lambda(x) = -\lambda x + \log(\lambda)$$

$$\ell_f = \sum_{i=1}^n -\lambda x_i + \log(\lambda) = -\lambda \sum x_i + n \log(\lambda)$$

$$\frac{\partial \ell_f}{\partial \lambda} = -\sum x_i + \frac{n}{\lambda} \Rightarrow \hat{\lambda} = \frac{1}{\bar{x}}$$

$$\text{Now } E_g \log f_\lambda(X) = E_g[-\lambda X + \log(\lambda)] = -\lambda E(X) + \log(\lambda)$$

$$\Rightarrow \text{maximized by } \lambda_0 = \frac{1}{E(X)}$$

$$\hat{\lambda} \rightarrow \lambda_0$$

If we believe f_λ is the true model, Asympt. Variance

$$\text{of } \hat{\lambda} \text{ is } \frac{F I_{\lambda_0}^{-1}}{n} = -\frac{1}{n} E \left(\frac{\partial^2}{\partial \lambda^2} \log f_\lambda(X) \right)_{\lambda_0}^{-1}$$

$$= -\frac{1}{n} \left(E \left[-\frac{1}{\lambda^2} \right] \right)^{-1} = \frac{\lambda_0^2}{n}$$

$$\text{so } \hat{\lambda} \sim N(\lambda_0, \frac{\lambda_0^2}{n})$$

in practice use

$$\hat{\lambda} - \lambda_0 \sim N(0, \frac{\hat{\lambda}^2}{n})$$

for testing etc

What if $f \neq g$

(9)

$$\Rightarrow J(\lambda_0) = E\left(\left(\frac{\partial}{\partial \lambda} \log f_\lambda(x)\right)^2 \middle| \lambda_0\right)$$

$$= E\left(-X + \frac{1}{\lambda_0}\right)^2 = E\left(X - \frac{1}{\lambda_0}\right)^2 = E(X - E(X))^2 = V(X)$$

/
depends on g .

$\Rightarrow$ robust variance

$$J\tilde{I}_{\lambda_0}^{-1} \tilde{V}_{\lambda_0} J\tilde{I}_{\lambda_0}^{-1} = \lambda^2 V(X) \lambda^2 = \lambda^4 V(X) = \text{var}$$

Example, $g \sim \chi_p^2$ $E(X) = p$ $V(X) = 2p$, $p = \text{positive integer}$

$$\Rightarrow \lambda = \frac{1}{p}$$

$$\Rightarrow \text{var} \sim \frac{2}{p^3} = 2\lambda^3 \quad \text{compared with } \lambda^2, \frac{1}{p^2}$$