

Model Selection

(7)

We used the LRT to compare models where a set of parameters were restricted. However, the LRT cannot be used to compare nonnested models, models with different dimensionality or family.

We consider models $f_j, j=1, \dots, J$ where

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

We cannot use $\ell(f_j, \hat{\theta}_j) = \log L_f(\hat{\theta}_j | \underline{x}) = \sum_{i=1}^n \log f_j(x_i; \hat{\theta}_j)$

Since model with max p_j 'wins' by construction

Focus on generalization - does the model, and the MLE under this model, work on future data?

Our criterion is

Choose f_j to maximize

$$\mathbb{E}_g [\log L_{f_j}(\theta_{f_j} | X)]$$

where we assume $X \sim^d g$ (unknown, true distribution)

Too many unknowns (don't know θ_{f_j}) $\Rightarrow$ replace with

$$\mathbb{E}_g [\log L_{f_j}(\hat{\theta}_{f_j}(\underline{x}) | X)] \equiv Q_{f_j}$$

expectation over both uncertainty in $\underline{x}$ and X

MLE under f_j using $\underline{x} = X_1^n \sim^d g$ (unknown)

generic, new future observation $X \sim^d g$

How to get a good and Computable estimate of Q_{f_j} ? ②

We observe $\tilde{x}$, can estimate $\hat{\theta}_{f_j}(\tilde{x})$ and the maximized likelihood $\frac{1}{n} \ell_{f_j}(\hat{\theta}_{f_j}(\tilde{x})) = \frac{1}{n} \sum_{i=1}^n \log f_j(x_i | \hat{\theta}_{f_j})$

• Note, $\frac{1}{n} \ell_{f_j}(\hat{\theta}_{f_j})$ is not a good estimate of Q_{f_j}

Since we use $\tilde{x}$ both for training (to obtain $\hat{\theta}$) and testing (to compute the likelihood).

$\frac{1}{n} \ell_{f_j}(\hat{\theta}_{f_j})$ is an overestimate of true Q_{f_j} .

Properties & notations (drop f_j for now and work with one f)

• $\hat{\theta}_f = \underset{\theta_f}{\operatorname{argmax}} \ell_f(\theta_f | \tilde{x})$

• $\frac{1}{n} \sum_{i=1}^n \log f(x_i | \theta_f) \rightarrow \mathbb{E}_g [\log f(X | \theta_f)]$ by WLLN

and $\max \text{ of LH} \rightarrow \max \text{ of RH}$

• Define θ_0 as maximizer of $\text{RH } \mathbb{E}_g [\log f(X | \theta_f)]$

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

• Note, θ_0 is not "all true" parameter unless $f=g$

Instead, θ_0 is the best choice of θ_f to match f

as closely as possible to g — best f -approximation of g

• $\hat{\theta}_f = \arg \max_{\theta_f} l_f(\theta_f | \tilde{X})$ or zero-crossing of the

(2)

Score $S_f(\theta_f) = \frac{\partial}{\partial \theta_f} \sum_{i=1}^n \log f(x_i | \theta_f) = \sum_{i=1}^n \frac{\partial}{\partial \theta_f} \log f(x_i | \theta_f)$

• $S_f(\theta_f) \approx S_f(\hat{\theta}_f) + \underbrace{\frac{\partial^2}{\partial \theta_f \partial \theta_f^T} \log L_f(\theta_f | \tilde{X})}_{\text{2nd order derivative}} \bigg|_{\theta^*} (\theta_f - \hat{\theta}_f)$

Approx near $\hat{\theta}_f$ by def. value between θ_f and $\hat{\theta}_f$

(For simplicity $\Rightarrow$ notation, drop subscript f on θ below)

So $S_f(\theta) = \frac{\partial^2}{\partial \theta \partial \theta^T} \log L_f(\theta | \tilde{X}) \bigg|_{\theta^*} (\theta - \hat{\theta})$

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

where $E(y_i) = E_y \left[\frac{\partial}{\partial \theta} \log f(x | \theta) \right] = \frac{\partial}{\partial \theta} E_y [\log f_{\theta}(x)]$
 $= 0$ at θ_0 since θ_0 maximizer of

At θ_0 , $V(y_i) = E(y_i^2) = E_y \left[\left(\frac{\partial}{\partial \theta} \log f_{\theta}(x) \right) \left(\frac{\partial}{\partial \theta} \log f_{\theta}(x) \right)^T \right]$
 $\equiv J(\theta_0)$

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

We have

(7)

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

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

we want distribution of this.

$$\cdot \frac{\partial^2}{\partial \theta \partial \theta^T} \log L_f(\theta | \tilde{x}) \Big|_{\theta_0} \approx \frac{\partial^2}{\partial \theta \partial \theta^T} \log L_f(\theta | \tilde{x}) \Big|_{\theta_0} \quad \text{for large } n$$

$$\cdot \frac{1}{n} \frac{\partial^2}{\partial \theta \partial \theta^T} \log L_f(\theta | \tilde{x}) \Big|_{\theta_0} \rightarrow E_g \left[\frac{\partial^2}{\partial \theta \partial \theta^T} \log f_\theta(X) \Big|_{\theta_0} \right] \text{ by WLLN}$$

$$\equiv -FI(\theta_0)$$

All together we have

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

$$\rightarrow^d N(0, J(\theta_0)) \quad \rightarrow^d \sqrt{-FI(\theta_0)}$$

$\Rightarrow$ Slutsky's Thm

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

Key result: CLT of MLE under wrong model

$$\text{if } f \neq g \quad FI(\theta_0)^{-1} = J(\theta_0) \Rightarrow \hat{\theta} \rightarrow^d N(\theta_0, (nFI(\theta_0))^{-1})$$

Comparing models

5

Back to original problem; $f_j, j=1, \dots, J$

• $\hat{\theta}_j = \text{maximizer of } \sum_{i=1}^n \log f_{\theta_j}(x_i) = l_{f_j}(\theta_j)$

• $\theta_{j0} = \text{maximizer of } E_g[\log f_{\theta_j}(X)]$

Want to choose f_j to maximize

$$Q = E_{\tilde{X}}[\log f_j(X | \hat{\theta}_j(\bar{X}))]$$

• We know $\sqrt{n}(\hat{\theta}_j - \theta_{j0}) \rightarrow^d N(0, FI_j^{-1}(\theta_{j0}) J_j(\theta_{j0}) FI_j^{-1}(\theta_{j0}))$

We start by examining the estimation variance
(dealing with randomness of $\tilde{x} \Rightarrow$ impact on $\hat{\theta}_j$)

We have $l_{f_j}(\theta_j) \approx l_{f_j}(\hat{\theta}_j) + \underbrace{\frac{\partial}{\partial \theta_j} l_{f_j}(\theta_j) |_{\hat{\theta}_j}}_{=0} (\theta_j - \hat{\theta}_j)$

$$+ \frac{1}{2} (\theta_j - \hat{\theta}_j)^T \left(\frac{\partial^2}{\partial \theta_j \partial \theta_j^T} l_{f_j}(\theta_j) |_{\theta_j^*} \right) (\theta_j - \hat{\theta}_j)$$

$$= l_{f_j}(\hat{\theta}_j) - \frac{1}{2} (\theta_j - \hat{\theta}_j)^T \boxed{**} (\theta_j - \hat{\theta}_j)$$

Where $\boxed{**} = -\frac{1}{n} \sum_{i=1}^n \frac{\partial^2}{\partial \theta_j \partial \theta_j^T} \log f_j(x_i | \theta_j) |_{\theta_j^*} \rightarrow^p E_g \left[\frac{\partial^2}{\partial \theta_j \partial \theta_j^T} \log f_j(X | \theta_j) |_{\theta_{j0}} \right]$
 $= FI_j^{-1}(\theta_{j0})$

So we have

⑥

$$l_{f_j}(\theta_j) \approx l_{f_j}(\hat{\theta}_j) - \frac{n}{2} (\theta_j - \hat{\theta}_j)^T \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0}) (\theta_j - \hat{\theta}_j) \quad (\text{A4*})$$

At θ_{j0} we have

$$\begin{aligned} & n(\hat{\theta}_j - \theta_{j0})^T \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0}) (\hat{\theta}_j - \theta_{j0}) \\ &= \left((n \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0}))^{1/2} (\hat{\theta}_j - \theta_{j0}) \right)^T \left((n \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0}))^{1/2} (\hat{\theta}_j - \theta_{j0}) \right) \\ &= \tilde{\theta}_j^T \tilde{\theta}_j \quad \tilde{\theta}_j \sim \text{standardized version of } \hat{\theta}_j \\ &= \sum_{p=1}^{p_j} \tilde{\theta}_{jp}^2 = \text{sum of diagonal elements of } \tilde{\theta}_j \tilde{\theta}_j^T \\ &= \text{Trace } \tilde{\theta}_j \tilde{\theta}_j^T \end{aligned}$$

We are interested in $E[\tilde{\theta}_j^T \hat{\theta}_j] = \text{Trace}(E[\tilde{\theta}_j \tilde{\theta}_j^T])$

$$= \text{Trace}(E[n \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0}) (\hat{\theta}_j - \theta_{j0}) (\hat{\theta}_j - \theta_{j0})^T])$$

$$= \text{Trace}(n \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0}) E[(\hat{\theta}_j - \theta_{j0}) (\hat{\theta}_j - \theta_{j0})^T])$$

$$\underbrace{Cov(\hat{\theta}_j)}_{\frac{1}{n}} = \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0})^{-1} \underbrace{J_j(\theta_{j0})}_{\frac{1}{n}} \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0})^{-1}$$

$$\Rightarrow \text{So } E[n(\hat{\theta}_j - \theta_{j0})^T \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0}) (\hat{\theta}_j - \theta_{j0})]$$

$$= \text{Trace} \left(n \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0}) \left(\mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0})^{-1} \underbrace{J_j(\theta_{j0})}_{\frac{1}{n}} \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0})^{-1} \right) \right) = \text{Trace}(J_j(\theta_{j0}) \mathcal{F}\bar{\mathcal{I}}_j(\theta_{j0})^{-1})$$

Back to test

(6)

$$l_{f_j}(\theta_j) \approx l_{f_j}(\hat{\theta}_j) - \frac{n}{2} (\hat{\theta}_j - \theta_{j0})^T F_{f_j}(\theta_{j0}) (\hat{\theta}_j - \theta_{j0})$$

and @ θ_{j0}

$$l_{f_j}(\theta_{j0}) \approx l_{f_j}(\hat{\theta}_j) - \frac{n}{2} (\hat{\theta}_j - \theta_{j0})^T F_{f_j}(\theta_{j0}) (\hat{\theta}_j - \theta_{j0})$$

$$\approx l_{f_j}(\hat{\theta}_j) - \frac{1}{2} \text{Trace} \left(J_{f_j}(\theta_{j0}) F_{f_j}(\theta_{j0}) \right)$$

[We used Taylor expansion of l_{f_j}

• CLT for $\hat{\theta}_j$ under f_j but g true model

• replacing 2nd order derivatives in approx with their expectations]

Take it all together;

$$l_{f_j}(\theta_{j0}) = \sum_{i=1}^n \log f_j(x_i | \theta_{j0}) \approx \sum_{i=1}^n \log f_j(x_i | \hat{\theta}_j) - \frac{1}{2} \text{Trace} \left(J_{f_j}(\theta_{j0}) F_{f_j}(\theta_{j0}) \right)$$

$$\Downarrow \mathbb{E}_{g, \tilde{x}} \Rightarrow n E[\log f_j(X | \theta_{j0})] = E \left[\sum_{i=1}^n \log f_j(x_i | \theta_{j0}) \right]$$

$$\approx E[l_{f_j}(\hat{\theta}_j)] - \frac{1}{2} \text{Trace} \left(J_{f_j}(\theta_{j0}) F_{f_j}(\theta_{j0}) \right)$$

Take a step back ; our goal is to maximize (wrt f_j)

(7)

$$Q_j = E_{\underset{\substack{\uparrow \\ \text{testing}}}{X}, \underset{\substack{\uparrow \\ \text{training}}}{\tilde{X}}} [\log f_j(X|\hat{\theta}_j; \tilde{X})]$$

We have so far done the training expectation, $E_{\tilde{X}}$.

To do testing, we expand

$$E_X [\log f_j(X|\theta_j)] \text{ near } \theta_{j0}$$

$$E_X [\log f_j(X|\theta_j)] = E [\log f_j(X|\theta_j)]$$

$$+ \underbrace{\frac{\partial}{\partial \theta_j} E [\log f_j(X|\theta_j)]}_{=0} \bigg|_{\theta_{j0}} (\theta_j - \theta_{j0})$$

$$+ \frac{1}{2} (\theta_j - \theta_{j0})^T \frac{\partial^2}{\partial \theta_j \partial \theta_j^T} E [\log f_j(X|\theta_j)] \bigg|_{\theta_{j0}} (\theta_j - \theta_{j0})$$

$$\approx E [\log f_j(X|\theta_{j0})] - \frac{1}{2} (\theta_j - \theta_{j0})^T E \left[\frac{\partial^2}{\partial \theta_j \partial \theta_j^T} \log f_j(X|\theta_{j0}) \bigg|_{\theta_{j0}} \right] (\theta_j - \theta_{j0})$$

$$= E [\log f_j(X|\theta_{j0})] - \frac{1}{2n} \left[n (\theta_j - \theta_{j0})^T \underbrace{E \left[\frac{\partial^2}{\partial \theta_j \partial \theta_j^T} \log f_j(X|\theta_{j0}) \bigg|_{\theta_{j0}} \right]}_{F_{T,j}(\theta_{j0})} (\theta_j - \theta_{j0}) \right]$$

(consider $\theta_j = \hat{\theta}_j$ especially;

⑧

$$E_x[\log f_j(X|\hat{\theta}_j)] \approx E[\log f_j(X|\theta_{j0})]$$

$$-\frac{1}{2n} [n(\hat{\theta}_j - \theta_{j0})^T F_{T,j}(\theta_{j0})(\hat{\theta}_j - \theta_{j0})]$$

Now take $E_{\tilde{x}}$ also $\Rightarrow$

$$E_{\tilde{x}}[E_x[\log f_j(X|\hat{\theta}_j)]] \approx E_x[\log f_j(X|\theta_{j0})]$$

(A)

$\parallel$
 Q_j

$$-\frac{1}{2n} \text{Trace}\{J_j(\theta_{j0}) F_{T,j}(\theta_{j0})^{-1}\}$$

from training error on page ⑥

From page ⑥ we have

$$n E_x[\log f_j(X|\theta_{j0})] \approx E_{\tilde{x}}[\log f_j(\hat{\theta}_j)] - \frac{1}{2} \text{Trace}\{J_j(\theta_{j0}) F_{T,j}(\theta_{j0})^{-1}\}$$

$$E_{\tilde{x}}[\sum_{i=1}^n \log f_j(X_i|\hat{\theta}_j(\tilde{x}))]$$

(B)

Combine (A) and (B) to see

$$n Q_j \approx E_{\tilde{x}}[\sum_{i=1}^n \log f_j(X_i|\hat{\theta}_j(\tilde{x}))] - \text{Trace}\{J_j(\theta_{j0}) F_{T,j}(\theta_{j0})^{-1}\}$$

(9)

We replace the average (expected) training error

$E_{\tilde{X}} \left[\sum_{i=1}^n \log f(x_i; \hat{\theta}_j(\tilde{x})) \right]$ by the observed

training error $\sum_{i=1}^n \log f_j(x_i; \hat{\theta}_j(\tilde{x})) = l_{f_j}(\hat{\theta}_j)$

$$\Rightarrow n Q_j \approx l_{f_j}(\hat{\theta}_j) - \text{Trace} \{ J_j(\theta_{j0}) F I_j(\theta_{j0})^{-1} \}$$

This depends on the unknown g
so cannot be computed.

BUT, assume f_j a good
approximation of g at θ_{j0} , then

$$J_j(\theta_{j0}) \approx F I_j(\theta_{j0})$$

$$\Rightarrow n Q_j \approx l_{f_j}(\hat{\theta}_j) - \text{Trace} \{ I_{p_j \times p_j} \}$$

$$= l_{f_j}(\hat{\theta}_j) - p_j$$

Usually we see this expressed as

$$ALL_j = -2n Q_j = \underbrace{-2l_{f_j}(\hat{\theta}_j)}_{\text{deviance under model } f_j} + \underbrace{2p_j}_{\text{penalty on model}}$$

deviance under
model f_j

penalty on model
 $p_j = 2 \times \# \text{ parameters}$

Strategy; • find $\hat{\theta}_j$: maximizer of l_{f_j} , for all f_j

(10)

• evaluate $ALL_j = -2l_{f_j}(\hat{\theta}_j) + 2p_j$

• choose f_j that minimizes ALL

Alternatives

• create a training and test set from original data $\tilde{X}$

• LODOCV, let each x_i take turns as X and $x_{i \neq i} = \tilde{X}$ training

Leave-one-out

Cross-validation

$$\hat{Q}_j = \frac{1}{n} \sum_{i=1}^n \log f_j(x_i | \hat{\theta}_j(x_{i \neq i}))$$

$$= LODOCV_j$$

• One can show $ALL_j \approx -2n LODOCV_j$

• can also try leave-k-out
cross-validation