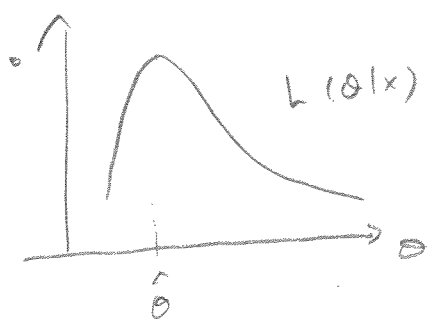


So far ; • likelihood function $L(\theta|x)$

①

- Shape of L captures uncertainty about θ in the sample data x
- To compare $L(\theta|x)$ for different values of θ in a data-scale independent fashion we use the likelihood ratio $LR = \frac{L(\theta_1|x)}{L(\theta_2|x)}$



$\Rightarrow$ summarize the likelihood function numerically?

(location of) maximum $\hat{\theta}_{MLE} = \underset{\theta}{\operatorname{argmax}} \{L(\theta|x) = 0\}$
 where $S(\theta|x) = \frac{\partial}{\partial \theta} \log L(\theta|x)$

if we can approximate $\log L(\theta|x)$ by a quadratic function we only need one more value; Curvature

$$@ \hat{\theta}, I(\hat{\theta}) = - \frac{\partial^2}{\partial \theta^2} \log L(\theta) \Big|_{\hat{\theta}}$$

$\Rightarrow$ Quadratic approximation (around $\hat{\theta}$)

$$\log L(\theta) \approx \log L(\hat{\theta}) + S(\hat{\theta})(\theta - \hat{\theta}) - \frac{1}{2} I(\hat{\theta})(\theta - \hat{\theta})^2$$

$$\Rightarrow \underbrace{\log \frac{L(\theta)}{L(\hat{\theta})}}_{\text{normalized log likelihood}} \approx -\frac{1}{2} I(\hat{\theta})(\theta - \hat{\theta})^2$$

normalized log likelihood

} (can plot and compare LHL & RHL to see if quad approx OK)

(2)

 $\Rightarrow$ follows that

$$S(\theta) \approx -I(\hat{\theta}) (\theta - \hat{\theta})$$

or

$$I(\hat{\theta})^{-1/2} S(\theta) \approx I(\hat{\theta})^{1/2} (\hat{\theta} - \theta)$$

(can show
(Wed) that

$$E_{\theta}(S(\theta)) = 0$$

$$V_{\theta}(S(\theta)) \approx I(\hat{\theta})$$

and approx
 $\sim N(0, 1)$

 $\Rightarrow$

$$\sim N(0, 1)$$

$$\Rightarrow \hat{\theta} \sim N(\theta, I(\theta)^{-1/2})$$

↓

testing

Confidence intervals

Note $W = 2 \log \frac{L(\hat{\theta})}{L(\theta)} \sim \chi^2_1$ under null $H_0: \theta$ (2)

- (came from know the sampling distribution of $\bar{x}$)

$\Rightarrow$ Interval estimate

$$\{ \theta : \frac{L(\theta)}{L(\hat{\theta})} > c \} = \{ \theta : 2 \log \frac{L(\hat{\theta})}{L(\theta)} < -2 \log c \}$$

$\sim \chi^2_1$

$\downarrow$
pick $-2 \log c = \chi^2_{1, (1-\alpha)}$
for $(1-\alpha)$ Conf. interval
or $c = e^{-\frac{1}{2} \chi^2_{1, (1-\alpha)}}$

$\Rightarrow$ Can also use for hypothesis testing

$H_0: \theta = \theta_0$

Reject if $2 \log \frac{L(\hat{\theta})}{L(\theta_0)} > \chi^2_{1, (1-\alpha)}$

Note if the quadratic approximation is poor

- no easy way to pick c .

For some special cases (x's and models for)

one can work out a probabilistic statement for

$\frac{L(\theta)}{L(\hat{\theta})}$ (see 2.6 in book)

Likelihood interval vs confidence interval

(3)

$$\left\{ \theta : \frac{L(\theta)}{L(\hat{\theta})} > c \right\} \quad (*)$$

Likelihood
interval

U plug in the quadratic approximation

$$\log \frac{L(\theta)}{L(\hat{\theta})} \approx -\frac{I(\hat{\theta})}{2} (\theta - \hat{\theta})^2 \quad \text{in } (*)$$

$$\Rightarrow \left\{ \theta : -\frac{I(\hat{\theta})}{2} (\theta - \hat{\theta})^2 > \log c \right\}$$

$$= \left\{ \theta : (\theta - \hat{\theta})^2 \leq I(\hat{\theta})^{-1} (-2 \log c) \right\}$$

$$= \left\{ \hat{\theta} \pm \sqrt{-2 \log c} \cdot I(\hat{\theta})^{-1/2} \right\} \quad (\text{Wald interval})$$

Remember, by quad approx $\hat{\theta} \sim N(\theta, \underbrace{I(\hat{\theta})^{-1}}_{V(\hat{\theta})})$

so pick $\sqrt{-2 \log c} = 1.96$ for a 95% confidence interval

Pro/Con w. Wilks vs Wald

Both are based on a quadratic approximation but Wilks (likelihood interval) allows for asymmetric intervals and tend to be better. Wald relies on approx $\frac{\hat{\theta} - \theta}{I(\hat{\theta})^{1/2}} \sim N(0, 1)$

whereas for Wilks we only need there to be some (unknown) $g(\theta)$ s.t. $g(\hat{\theta}) - g(\theta) / se(g(\hat{\theta})) \sim N(0, 1) \Rightarrow 2 \log \frac{L(\hat{\theta})}{L(\theta)} \sim \chi^2_1$

Quick look @ multi-parameter case...

- In scalar case — can plot $L(\theta)$
- In multi-parameter case — difficult to graph directly
 - issues to consider — can likelihood be factorized? how related are the components of θ ?

Notation & Terminology $\underline{\theta} = \{\theta_1, \dots, \theta_p\}$

$$\begin{aligned} \bullet \quad \underset{p \times 1}{S(\underline{\theta})} &= \nabla_{\underline{\theta}} \log L(\underline{\theta}) = \text{vector of 1st order derivatives} \\ &= \frac{\partial}{\partial \theta_j} \log L(\underline{\theta}) \end{aligned}$$

$$\bullet \quad \hat{\underline{\theta}}_{MLE} = \{ \underline{\theta} : S(\underline{\theta}) = \underline{0} \}$$

$$\begin{aligned} \bullet \quad \underset{\substack{p \times p \\ \text{(Observed)}}}{\text{Fisher information matrix}} \quad \mathcal{I}_{ij}(\hat{\underline{\theta}}) &= - \frac{\partial^2 \log L(\underline{\theta})}{\partial \theta_i \partial \theta_j} \bigg|_{\hat{\underline{\theta}}} \end{aligned}$$

because now need to care about curvature in high-dim space.

(can still try to approximate $L(\underline{\theta})$ by a quadratic shape (2nd order approx))

$$\log L(\underline{\theta}) \approx \log L(\hat{\underline{\theta}}) + \underset{1 \times 1}{S(\hat{\underline{\theta}})'} \underset{1 \times 1}{(\underline{\theta} - \hat{\underline{\theta}})} - \frac{1}{2} \underset{1 \times p}{(\underline{\theta} - \hat{\underline{\theta}})'} \underset{p \times p}{\mathcal{I}(\hat{\underline{\theta}})} \underset{p \times 1}{(\underline{\theta} - \hat{\underline{\theta}})}$$

$= 0$

Would be nice if $I(\hat{\theta})$ was diagonal

(5)

- separate uncertainty about components.

$\Rightarrow$ more later

Working w. multiparameter models in practise

$\Rightarrow$ In complex data situations, not so

easy to solve score equation $S(\hat{\theta}) = 0$

for all components of data simultaneously.

parameters of interest
+

'nuisance' parameters

$\{x, \mu, \sigma^2$

(more later)

complex dependencies

between components of θ or restrictions

$\rightarrow$ constrained optimization, unstable solution

Ex. Mixture model

$$X \sim \sum_u \pi_u N(Z' \theta_u, \Sigma_u)$$

some restricted to 0

One way to approach this is to

use profile likelihood

Definition

6

θ - parameter of interest

η - other parameter(s)

$L(\theta, \eta)$ (full) joint likelihood

$L(\theta) \equiv \max_{\eta} L(\theta, \eta)$ is called the profile likelihood for θ

max is taken for a fixed value of θ .
(unknown in general)

$\Rightarrow$ usually $\Rightarrow$ MLE of η for fixed θ is a function of θ

$$\Rightarrow \max_{\eta} L(\theta, \eta) = \eta(\theta)$$

$$\Rightarrow L(\theta) = L(\theta, \eta(\theta)) \Rightarrow \text{now maximize wrt } \theta$$

$\neq$ Integrating out η in joint distribution
is a Bayesian sense

$$p(x) \pi(\theta) = \int_{\eta} p(x|\theta, \eta) \pi(\theta, \eta) d\eta$$

Example normal-mean model

$$L(\mu, \sigma^2) = (\pi \sigma^2)^{-n/2} \exp \left\{ -\frac{1}{2\sigma^2} \sum (x_i - \mu)^2 \right\}$$

For fixed $\mu \Rightarrow \arg \max_{\sigma^2} L(\mu, \sigma^2) = \hat{\sigma}_{\mu}^2 = \frac{1}{n} \sum (x_i - \mu)^2$
 fixed unknown

$$\Rightarrow L(\mu) = (\text{constant} \times (\hat{\sigma}_{\mu}^2)^{-n/2}$$

1) Note : not the same as first $\max_{\mu, \sigma^2} L(\mu, \sigma^2) = L(\hat{\mu}, \hat{\sigma}^2)$

and plugging in $\hat{\sigma}^2$ into $L(\mu, \sigma^2)$

2) Note : $L(\hat{\sigma}^2) = \max_{\mu} L(\mu, \sigma^2)$ easy $\Rightarrow$ same $\hat{\mu}_{\hat{\sigma}^2} = \hat{\mu}_{MLE} = \bar{x}$

$$= (\text{constant} \cdot (\hat{\sigma}^2)^{-n/2} \exp \left\{ -\frac{1}{2\hat{\sigma}^2} \sum (x_i - \bar{x})^2 \right\}$$

$$n \hat{\sigma}^2$$

- Curse of dimensionality - hard to do inference in high-dim case.
- $L(\theta)$ can have complex shape and small values so the normalized likelihood not so unusual.
- So - picking reasonable cutoffs for $\frac{L(\theta)}{L(\hat{\theta})} \Rightarrow$ Interval estimates etc, requires that we take dimensionality of the problem into account.

Recap single case:

Wilks stat (good approx)
large sample

$$W = 2 \log \frac{L(\hat{\theta})}{L(\theta)} \sim \chi^2_1$$

$\Rightarrow$ What is a reasonable cutoff c for $\frac{L(\theta)}{L(\hat{\theta})} > c$?

(so $L(\theta)$ not too small comp $L(\hat{\theta})$)

$$\Rightarrow P\left(\frac{L(\theta)}{L(\hat{\theta})} > c\right) = P\left(\underbrace{2 \log \frac{L(\hat{\theta})}{L(\theta)}}_{\chi^2_1} < -2 \log c\right) = 1 - \alpha$$

i) we pick

$$-2 \log c = \chi^2_1(1 - \alpha)$$

$$\left\{ \theta : \frac{L(\theta)}{L(\hat{\theta})} > e^{-\frac{1}{2} \chi^2_1(1 - \alpha)} \right\}$$

$$\begin{aligned} & H_0: \theta = \theta_0 \\ & \text{Test/ing: reject } H_0 \text{ if } \\ & 2 \log \frac{L(\hat{\theta})}{L(\theta_0)} > \chi^2_1(1 - \alpha) \end{aligned}$$

Multparameter case

$$W = 2 \log \frac{L(\hat{\theta})}{L(\theta)} \sim \chi^2_p \quad (\text{approx})$$

$$\Rightarrow \left\{ \theta : \frac{L(\theta)}{L(\hat{\theta})} > \exp\left\{-\frac{1}{2} \chi^2_p(1 - \alpha)\right\} \right\}$$

Usually used for testing

(9)

$$H_0: \theta = \theta_0 \quad W = 2 \log \frac{L(\hat{\theta})}{L(\theta_0)} \sim \chi_p^2 \text{ under null}$$

$\Rightarrow$ Reject H_0 if $W > \chi_p^2(1-\alpha)$

$$\text{pvalue} \quad P(\chi_p^2 > \underbrace{2 \log \frac{L(\theta_0)}{L(\hat{\theta})}}_{\text{observed LR}}) \quad (10)$$

observed LR

Notice how cutoff or critical value of test $\chi_p^2(1-\alpha)$

depends on p . So a LR of 25% means something

different for $p=2$ or 10 .

That is, we cannot use a fixed LR comparison but have to take dimensionality of problem into account

In fact, for fixed $LR = \frac{L(\theta_0)}{L(\hat{\theta})}$ one can show that

the pvalue (10) $\rightarrow 1$ as $p \rightarrow \infty$ so with a fixed LR cutoff we would reject too often.

Heuristic derivation for AIC (model selection) (15)

Want a good model, but as simple as possible

Example, null model $\theta_0 = 0$ (parameters = 0 for some subset of parameters)

(like in regression modeling)

'Best' model to match data is the MLE $\hat{\theta}$, but that can be an overfit with many spurious effects estimated (all components of $\hat{\theta}$ are nonzero)

• Let's say $L(\theta) \approx L(\theta_1)L(\theta_2) \dots L(\theta_p)$ (factorizes parameter components)

• In scalar case we view $\frac{L(\theta_0)}{L(\hat{\theta})} \leq c$ as an indication that θ_0 is a poor fit (for some value c , e.g. 10%)

• Let's now say we want a similar support for all components of θ to prefer model θ_0

$$\frac{L(\theta_{01})}{L(\hat{\theta}_1)} > c, \frac{L(\theta_{02})}{L(\hat{\theta}_2)} > c \quad \text{etc}$$

• Or $\frac{L(\theta_0)}{L(\hat{\theta})} > c^p$ to "accept" θ_0 model or

$\frac{L(\theta_0)}{L(\hat{\theta})} \leq c^p$ to reject it (rejecting all!)

That is how much $(2p)$ we think $-2\log L(\theta)$ (12)
can be reduced compared to fixed $-2\log L(\theta^*)$ just
by chance

Strategy : Minimize $-2 \times \text{log likelihood} + \text{penalty } 2p$

$\Rightarrow$ Winning model finds a good balance between
fit ($-2\log \text{like}$) and complexity (# of parameters)

The penalty $2p$ is more aggressive than $\chi^2_p(1-\alpha)$
cutoff in Wilks so AIC tends to pick smaller
models (further from null) more often than Wilks