

Mathematical Statistics II

STA2212H S LEC9101

Week 2

January 17 2023

Two-thirds of world's glaciers expected to disappear by end of the century, study in Science journal says

SETH BORENSTEIN

But if the world can limit future warming to just a few more tenths of a degree and fulfill international goals – technically possible but unlikely according to many scientists – then slightly less than half the globe's glaciers will disappear, said the same study. Mostly small but wellknown glaciers are marching to extinction, study authors said.

In an also unlikely worst-case scenario of several degrees of warming, 83 per cent of the world's glaciers would likely disappear by the year 2100, study authors said.

The study, published Thursday in the journal Science, examined all of the globe's 215,000 landbased glaciers – not counting those on ice sheets in Greenland and Antarctica – in a more comprehensive way than past studies. Scientists



Tourists hike to visit the Nigardsbreen glacier in Jostedal, Norway, last August. Scientists project the planet will lose between 38.7 trillion and 64.4 trillion tonnes of glacial ice by the end of the century.

1. Recap
2. Nonparametric Likelihood [MS 5.6](#)
3. Profile Likelihood
4. Bayesian Estimation [MS 5.8](#)

Upcoming seminars of interest

- [January 23 11.00 –12.00 Jevin West Details](#)
- “The Art of Skepticism in a Data-Driven World”
- 140 St. George St., 4th floor





HOME

SYLLABUS

BOOK

VIDEOS

TOOLS

CASE STUDIES

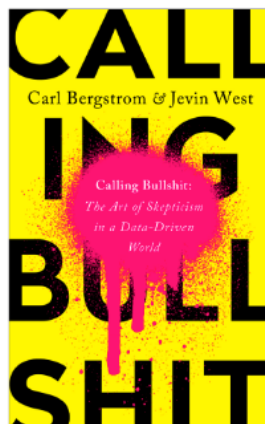
FAQ

PRESS

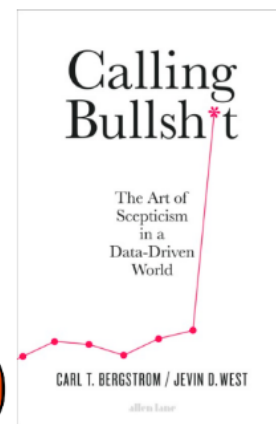
CONTACT

Calling Bullshit

Data Reasoning in a Digital World



Penguin
Random
House



Now available! *Calling Bullshit: The Art of Skepticism in a Data-Driven World*, by Carl Bergstrom and Jevin West. [Available here.](#)

Recap

- data $x_1, \dots, x_n$ independent observations; model $f(\mathbf{x}; \theta) = \prod_{i=1}^n f(x_i; \theta)$, $\theta \in \mathbb{R}$
- limit theorem $\sqrt{n}(\hat{\theta} - \theta) \xrightarrow{d} N(0, I_1^{-1}(\theta)) = E\{\ell'(\theta; x_i)^2\}$ $\hat{\theta}_n \rightarrow \theta$
- approximation $\hat{\theta} \sim N\{\theta, I^{-1}(\hat{\theta})\}$, or $\hat{\theta} \sim N\{\theta, j^{-1}(\hat{\theta})\}$ $I(\theta) = nI_1(\theta)$, $j(\theta) = -\ell''(\theta; \mathbf{x})$
observed value

Recap

- data $x_1, \dots, x_n$ independent observations; model $f(\mathbf{x}; \theta) = \prod f(x_i; \theta)$, $\theta \in \mathbb{R}$
- limit theorem $\sqrt{n}(\hat{\theta} - \theta) \xrightarrow{d} N(0, I_1^{-1}(\theta))$
- approximation $\hat{\theta} \sim N\{\theta, I^{-1}(\hat{\theta})\}$, or $\hat{\theta} \sim N\{\theta, J^{-1}(\hat{\theta})\}$ $l(\theta) = n l_1(\theta), \quad J(\theta) = -\ell''(\theta; \mathbf{x})$

- data $x_1, \dots, x_n$ independent observations; model $f(\mathbf{x}; \theta) = \prod f(x_i; \theta)$, $\theta \in \mathbb{R}^p$
- limit theorem $\sqrt{n}\{I(\theta)\}^{1/2}(\hat{\theta} - \theta) \xrightarrow{d} N_p(\mathbf{0}, I_d)$ $\sqrt{n}(\hat{\theta} - \theta) \xrightarrow{d} N_p(\mathbf{0}, I_1(\theta))$
 $\uparrow$
 $p \times p$
matrix
- approximation $\hat{\theta} \sim N_p\{\theta, I^{-1}(\hat{\theta})\}$, or $\hat{\theta} \sim N_p\{\theta, J^{-1}(\hat{\theta})\}$ $\nabla \frac{\partial^2 \ell(\hat{\theta}; \mathbf{x})}{\partial \theta \partial \theta^T}$

Recap

- data $x_1, \dots, x_n$ independent observations; model $f(\mathbf{x}; \theta) = \prod f(x_i; \theta)$, $\theta \in \mathbb{R}$
- limit theorem $\sqrt{n}(\hat{\theta} - \theta) \xrightarrow{d} N(0, I_1^{-1}(\theta))$
- approximation $\hat{\theta} \sim N\{\theta, I^{-1}(\hat{\theta})\}$, or $\hat{\theta} \sim N\{\theta, J^{-1}(\hat{\theta})\}$ $l(\theta) = nl_1(\theta)$, $J(\theta) = -\ell''(\theta; \mathbf{x})$

- data $x_1, \dots, x_n$ independent observations; model $f(\mathbf{x}; \theta) = \prod f(x_i; \theta)$, $\theta \in \mathbb{R}^p$
- limit theorem $\sqrt{n}\{l(\theta)\}^{1/2}(\hat{\theta} - \theta) \xrightarrow{d} N_p(\mathbf{0}, I_d)$
- approximation $\hat{\theta} \sim N_p\{\theta, I^{-1}(\hat{\theta})\}$, or $\hat{\theta} \sim N_p\{\theta, J^{-1}(\hat{\theta})\}$

cdf $\hat{F}_n(\cdot)$

- data $x_1, \dots, x_n$ independent observations id. dist^d obs = with cdf $F(\cdot)$
- true model $F(\mathbf{x}) = \prod F(x_i)$, $\theta \in \mathbb{R}^p$ assumed model $\ell(\theta; \mathbf{x})$, $\ell'(\hat{\theta}; \mathbf{x}) = \mathbf{0}$ Δ

- limit theorem $\sqrt{n}\{\hat{\theta} - \theta(F)\} \xrightarrow{d} N\{\mathbf{0}, \underbrace{J^{-1}(F)I(F)J^{-1}(F)}_P\}$

$\theta(F), I(F), J(F)$

$$J(F) = \mathbb{E} \left\{ \ell''(\theta_F; \mathbf{x}) \right\}$$

$$I(F) = \mathbb{E} \left\{ \ell'(\theta_F; \mathbf{x})^2 \right\}$$

$$\mathbb{E}_F \{ \ell'(\theta_F; x_i) \} = 0$$

... Recap

- proof requires many smoothness conditions on underlying model
- **i.i.d.** can often be weakened to independent (not i.d.) observations, or even dependent **AR(1), AR(p), GLMs (regression)** need WLLN and CLT
- MS Theorem 5.3, p.253 has a careful proof for $\theta \in \mathbb{R}$

A1 - A6

B1 - B6

see also MSI, Nov 29, likelihood handout

- key step is

$$\ell'(\hat{\theta}; \underline{X}) = 0$$

$$= \dots \text{TS}$$



$$\sqrt{n}(\hat{\theta} - \theta) = \frac{+n^{-1/2} \sum_{i=1}^n \ell'(X_i; \theta) \xleftarrow{\text{CLT}}}{-n^{-1} \sum_{i=1}^n \ell''(X_i; \theta) + (\hat{\theta} - \theta) (2n)^{-1} \sum_{i=1}^n \ell'''(X_i; \theta^*)}$$

$\downarrow$ $\downarrow$ $\downarrow$ $\downarrow$ $\downarrow$
 $E_{\theta} \ell''(X_i; \theta)$ 0 R_n $E_{\theta} \ell'''(X_i; \theta^*)$ $E_{\theta} \ell'''(X_i; \theta^*)$
 WLLN

$(\theta^* - \theta) < |\hat{\theta} - \theta|$
 $\hat{\theta} \xrightarrow{P} \theta$
 $\Rightarrow \theta^* \xrightarrow{P} \theta$

... Recap

- proof requires many smoothness conditions on underlying model
- **i.i.d.** can often be weakened to independent (not i.d.) observations, or even dependent
- MS Theorem 5.3, p.253 has a careful proof for $\theta \in \mathbb{R}$

need WLLN and CLT

see also MSI, Nov 29, likelihood handout

- key step is

$$\sqrt{n}(\hat{\theta} - \theta) = \frac{-n^{-1/2} \sum_{i=1}^n \ell'(X_i; \theta)}{n^{-1} \sum_{i=1}^n \ell''(X_i; \theta) + (\hat{\theta} - \theta)(2n)^{-1} \sum_{i=1}^n \ell'''(X_i; \theta^*)}$$

- vector version is

$$\sqrt{n} \sum_{k=1}^p (\hat{\theta}_k - \theta_k) \left\{ \underbrace{n^{-1} \ell''_{jk}(\hat{\theta})}_{\frac{\partial^2 \ell}{\partial \theta_j \partial \theta_k}} + (2n)^{-1} \sum_{l=1}^p (\hat{\theta}_l - \theta_l) \underbrace{\ell'''_{jkl}(\theta^*)}_{\frac{\partial^3 \ell}{\partial \theta_j \partial \theta_k \partial \theta_l}} \right\} = \underbrace{-n^{-1/2} \ell'_j(\theta)}_{j=1, \dots, p},$$

$$\ell'(\hat{\theta}) = \underline{0} \approx \ell'(\theta) + \dots$$

$$\theta \in \mathbb{R}^p$$

- proof of consistency (Thms 5.1,2) uses WLLN applied to

θ fixed (true)
 t varying

$$\phi_n(t) = \frac{1}{n} \sum_{i=1}^n \log \frac{f(X_i; t)}{f(X_i; \theta)}$$

$$\phi(t) = E_{\theta} \log \frac{f(X_i; t)}{f(X_i; \theta)} \equiv -K(f_t : f_{\theta})$$

$$\phi_n(t) \rightarrow \phi(t) \text{ for any } t \in \mathbb{R}^p$$

$$-K(f_t : f_{\theta}) = 0 \text{ iff } t = \theta$$

$$\text{o.w. } -K < 0$$

$$K(\text{true})$$

$$\sup_t \phi_n(t) \neq \text{determines } \hat{\theta}$$

$$M_n(\theta) = \frac{1}{n} \sum_{i=1}^n \log \frac{f(X_i; \theta)}{f(X_i; \theta_{\text{true}})}$$

$$E_{\theta_{\text{true}}} \log \frac{f(X_i; \theta)}{f(X_i; \theta_{\text{true}})} \equiv -D(\theta_{\text{true}}, \theta)$$

$$\cancel{A1} = A6$$

$$B1 - B6$$

$$\theta^*$$

As

$$\sup_t \phi(t) = \phi(\theta)$$

As

$$\sup_{\theta \in \Theta} |M_n(\theta)| < C$$

$$\Theta \subseteq \mathbb{R}$$

$$\sup_{\{K, \sigma^2 \in \mathbb{R}^n \times \mathbb{R}^+\}} \left| \frac{\partial \log f(x; \theta)}{\partial \theta} \right|$$

$$\hat{\mu}_{i,n} = \frac{1}{2} (X_i + Y_i) \quad i=1, \dots, n$$

$\hat{\mu}_{i,n} \rightarrow \mu_i$ b.c. as $n \rightarrow \infty$ $\hat{\mu}_{i,n}$ doesn't change

- proof of consistency (Thms 5.1,2) uses WLLN applied to

$$\phi_n(t) = \frac{1}{n} \sum_{i=1}^n \log \frac{f(X_i; t)}{f(X_i; \theta)}$$

$$M_n(\theta) = \frac{1}{n} \sum_{i=1}^n \log \frac{f(X_i; \theta)}{f(X_i; \theta_{true})}$$

$$\phi(t) = E_{\theta} \log \frac{f(X_i; t)}{f(X_i; \theta)} \equiv -K(f_t : f_{\theta})$$

$$E_{\theta_{true}} \log \frac{f(X_i; \theta)}{f(X_i; \theta_{true})} \equiv -D(\theta_{true}, \theta)$$

$\Leftarrow$

true

- by Jensen's $\phi(t)$ maximized at θ , which suggests $\hat{\theta} \rightarrow \theta$
- need sup condition, see Thm 5.1 (a)

but functions are tricky

p fixed
n → ∞

- proof of consistency (Thms 5.1,2) uses WLLN applied to

$$\phi_n(t) = \frac{1}{n} \sum_{i=1}^n \log \frac{f(X_i; t)}{f(X_i; \theta)}$$

$$M_n(\theta) = \frac{1}{n} \sum_{i=1}^n \log \frac{f(X_i; \theta)}{f(X_i; \theta_{true})}$$

$$\phi(t) = E_{\theta} \log \frac{f(X_i; t)}{f(X_i; \theta)} \equiv -K(f_t : f_{\theta})$$

$$E_{\theta_{true}} \log \frac{f(X_i; \theta)}{f(X_i; \theta_{true})} \equiv -D(\theta_{true}, \theta)$$

$$\int \log \left\{ \frac{f(x)}{f_0(x)} \right\} f_0(x) dx$$

but functions are tricky

- by Jensen's $\phi(t)$ maximized at θ , which suggests $\hat{\theta} \rightarrow \theta$
- need sup condition, see Thm 5.1 (a) *also AoS Chapter 9*

base density

- $K(f : f_0)$ Kullback-Leibler divergence measures 'closeness' of densities f and f_0

$$E_0 \{ \log \{ f_0(X) / f(X) \} \}$$

- maximum likelihood estimator minimizes K-L divergence between empirical cdf and model

$$E_{F_n} \log \{ dF_n(\mathbf{x}) / f_{\theta}(\mathbf{x}) \}$$

- sample $x_1, \dots, x_n$ independent, identically distributed, with cdf F

(assume F has density)

no parametric model assumed

- if I use this
- likelihood function $L(\mathbf{f}) = \prod_{i=1}^n f(x_i)$ $\times$

- assume solution puts mass only at $x_1, \dots, x_n$

- log-likelihood function $\ell(\mathbf{p}) = \sum_{i=1}^n \log(p_i)$

$$\sum p_i = 1 \quad 0 \leq p_i \leq 1$$

(ntbc)

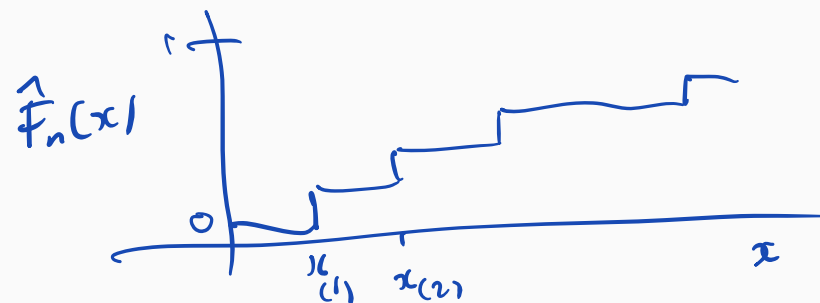
$$\max_{\mathbf{p}} \sum_{i=1}^n \log(p_i) \quad ; \quad \text{s.t.} \quad \sum p_i = 1$$

$$0 < p_i < 1$$

$$p_i = \frac{1}{n} \quad i=1, \dots, n$$

$$\hat{F}_n(x) = \frac{1}{n} \sum 1\{x_i \leq x\}$$

cdf



- sample $x_1, \dots, x_n$ independent, identically distributed, with cdf F

no parametric model assumed

- likelihood function $L(F) = \prod f(x_i)$

$$f_x(p) = \ell(p) + \lambda (\sum p_i - 1)$$

$$\frac{\partial}{\partial \lambda} \quad \frac{\partial}{\partial p_1} \quad \dots \quad \frac{\partial}{\partial p_n}$$

- assume solution puts mass only at $x_1, \dots, x_n$

- log-likelihood function $\ell(p) = \sum_{i=1}^n \log(p_i)$

- maximized at $p_i = 1/n, i = 1, \dots, n$

Lagrange

- gives empirical cdf

$$\hat{F}_n(x) = \frac{1}{n} \sum_{i=1}^n 1(X_i \leq x)$$

$\hat{F}_n(\cdot)$ is "MLE" of $F(\cdot)$

$$\hat{f}_n(x) = \frac{1}{nh} \sum_{i=1}^n K\left(\frac{x_i - x}{h}\right) \quad ?? \text{ nt bc}$$



Multi-parameter example: logistic regression

```
Boston$crim2 <- Boston$crim > median(Boston$crim) # define binary response
Boston.glm <- glm(crim2 ~ . - crim, family = binomial,
                  data = Boston) #fit logistic regression
summary(Boston.glm)
```

Coefficients:

	Estimate	Std. Error	z value	Pr(> z)
(Intercept)	-34.103704	6.530014	-5.223	1.76e-07 ***
zn	-0.079918	0.033731	-2.369	0.01782 *
indus	-0.059389	0.043722	-1.358	0.17436
chas	0.785327	0.728930	1.077	0.28132
nox	48.523782	7.396497	6.560	5.37e-11 ***
rm	-0.425596	0.701104	-0.607	0.54383
age	0.022172	0.012221	1.814	0.06963 .

Wald st.
 $\hat{\theta}_j \sim N(\theta_j, \dots)$

$$\frac{\partial \ell(\hat{\theta})}{\partial \theta_j} = 0$$

$$\left[\begin{array}{c} \frac{\partial \ell(\hat{\theta})}{\partial \theta_1} = 0 \\ \vdots \\ \frac{\partial \ell(\hat{\theta})}{\partial \theta_p} = 0 \end{array} \right]$$

... Example: logistic regression

```
Boston.glm <- glm(crim2 ~ . - crim, family = binomial,  
                  data = Boston) #fit logistic regression
```

```
confint(Boston.glm)
```

Waiting for profiling to be done...

2.5 % 97.5 %

(Intercept) -47.480389822 -21.699753794

zn -0.152359922 -0.020567540

indus -0.149113408 0.024168460

chas -0.646429219 2.233443233

nox 34.967619055 64.088411260

rm -1.811639107 0.950196261

age -0.001231256 0.046865843

dis 0.280762523 1.140619391

rad 0.376833861 0.975898274

tax -0.012038221 -0.001324887

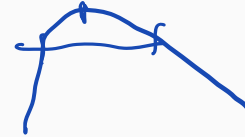
$$\hat{\beta}_j \pm t.96(\hat{se}_j)$$



95% CI

for β_j

$\ell_p(4)$



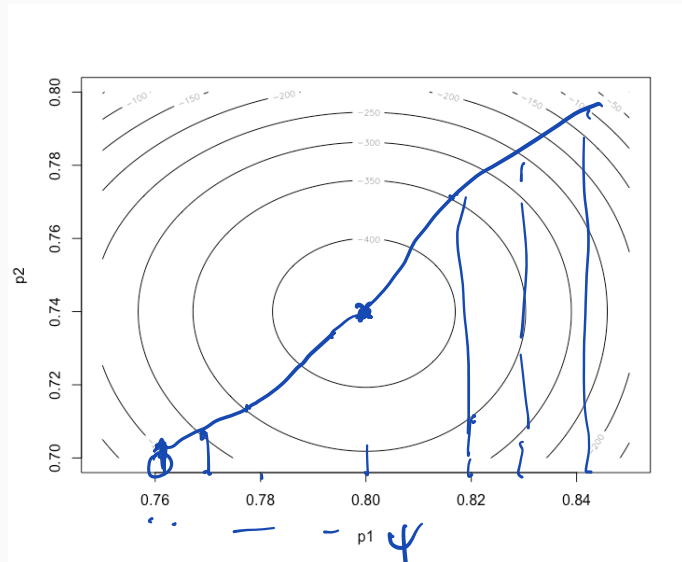
$$-0.08 \pm 0.06$$

$$(-0.02, 0.14)$$

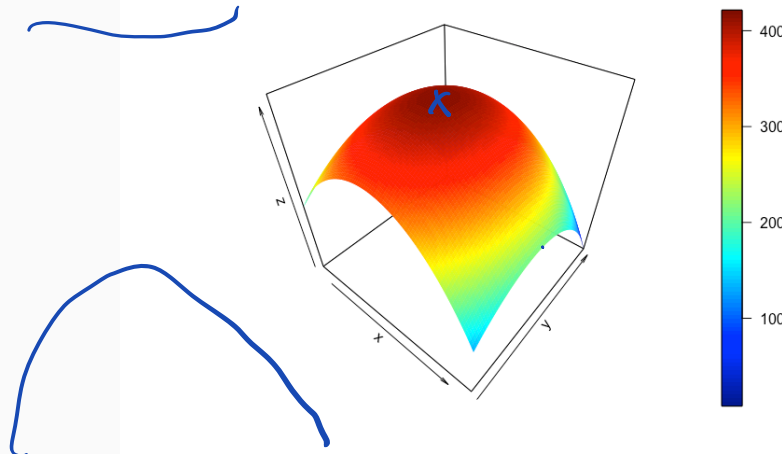
Prof.

$$(-0.14, -0.02)$$

$$\hat{\beta}_j \pm \dots$$



$$\hat{p}_1 = \frac{160}{200} \quad \hat{p}_2 = \frac{180}{200}$$



$$\frac{\partial \ell(\hat{\theta})}{\partial \theta} = 0$$

$p \times 1 \quad p \times 1$

$Y_1 \sim \text{Binom}(n_1, p_1)$, $Y_2 \sim \text{Binom}(n_2, p_2)$, independently
observed values $y_1 = 160$, $n_1 = 200$, $y_2 = 180$, $n_2 = 200$

$$Y_1 \sim B(200, p_1)$$

$$\frac{\partial \ell}{\partial p_1} = 0 \quad \dots \quad \hat{p}_1(p_2)$$

$$Y_2 \sim B(200, p_2)$$

$$\ell(p_1, p_2) = \log \binom{200}{y_1} p_1^{y_1} (1-p_1)^{200-y_1} + \log \binom{200}{y_2} p_2^{y_2} (1-p_2)^{200-y_2}$$

$$\left\{ \begin{array}{l} \frac{\partial l(\theta_1^i, \theta_2^i; x)}{\partial \theta_1} = 0 \\ \frac{\partial l(\theta_1, \theta_2^i; x)}{\partial \theta_2} = 0 \end{array} \right.$$

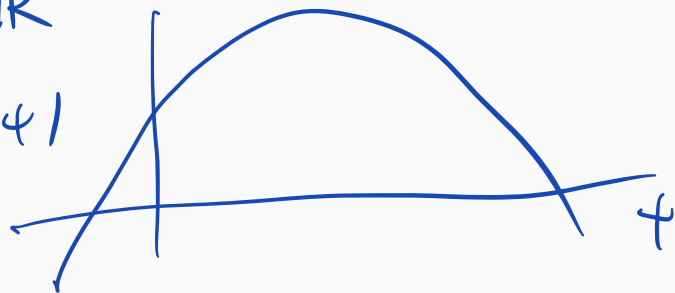
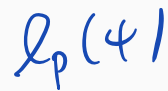
Profile likelihood function

often $\underline{\theta} = (\psi, \underline{\lambda})$

- ψ - parameter of interest (σ^2)
- $\underline{\lambda}$ - nuisance parameters ($\mu, \dots, \mu$)

define $\ell_p(\psi) = \ell(\psi, \hat{\lambda}_\psi)$ $\frac{\partial \ell}{\partial \lambda}(\psi, \hat{\lambda}_\psi) = 0$

if $\psi \in \mathbb{R}$



concentrated liq.

... Profile likelihood function

$$l_p(\psi) = \underline{l(\psi, \hat{\lambda}_\psi)}$$

$$\underline{l'_p(\tilde{\psi}) = 0}$$

$$\tilde{\psi} \equiv \hat{\psi}$$

$$0 = l'_p(\tilde{\psi}) \Big|_{\psi=\tilde{\psi}} = \underbrace{\frac{\partial}{\partial \psi} \{ l(\psi, \hat{\lambda}_\psi) \}}_{\psi=\tilde{\psi}} + \cancel{\frac{\partial}{\partial \lambda} l(\psi, \hat{\lambda}_\psi) \frac{d\hat{\lambda}_\psi}{d\psi}}_{\text{by def of } \hat{\lambda}_\psi}$$

$$\left[\begin{array}{l} \frac{\partial l}{\partial \lambda}(\tilde{\psi}, \hat{\lambda}_{\tilde{\psi}}) = 0 \\ \frac{\partial l}{\partial \psi}(\tilde{\psi}, \hat{\lambda}_{\tilde{\psi}}) = 0 \end{array} \right] \iff \left(\begin{array}{c} \frac{\partial l(\psi, \lambda)}{\partial \psi} \\ \frac{\partial l(\psi, \lambda)}{\partial \lambda} \end{array} \right) \Big|_{(\tilde{\psi}, \hat{\lambda})} = 0$$

$\Rightarrow$ Profile $\log$ -lik. work like $\log$ -lik
 (asymptotically)

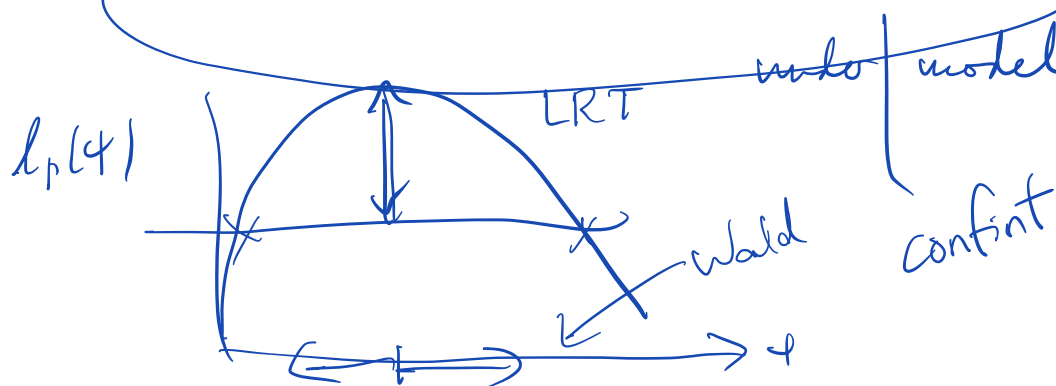
theorems
 p fixed
 $n \rightarrow \infty$

(i) $l'_p(\hat{\psi}) = 0$ $\hat{\psi}$ m.l.e.

(ii) $\widehat{\text{a.var.}}(\hat{\psi}) = -l''_p(\hat{\psi})$ obs'd Fisher info in profile

(approxⁿ may be poor, p large)

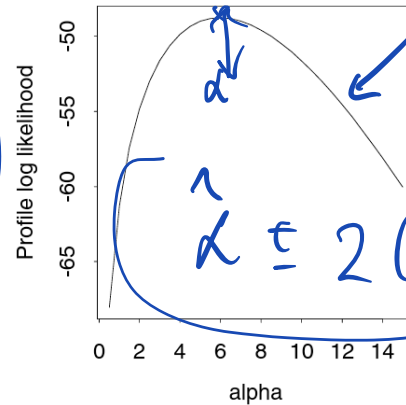
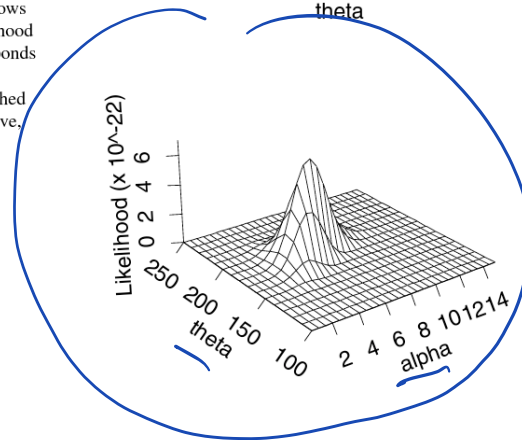
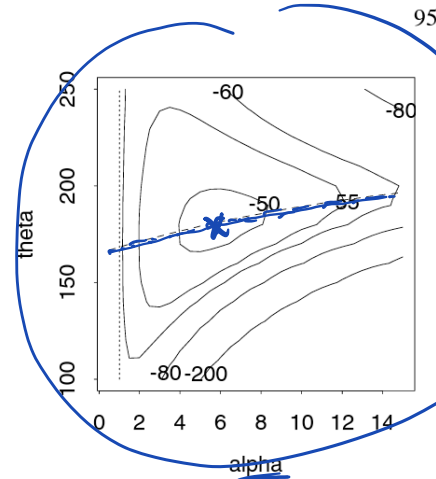
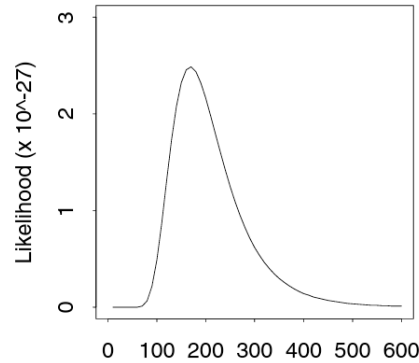
(iii) $2\{l_p(\hat{\psi}) - l_p(\psi)\} \sim \chi^2_1$



... Profile likelihood function

4.1 · Likelihood

Figure 4.1 Likelihoods for the spring failure data at stress 950 N/mm². The upper left panel is the likelihood for the exponential model, and below it is a perspective plot of the likelihood for the Weibull model. The upper right panel shows contours of the log likelihood for the Weibull model; the exponential likelihood is obtained by setting $\alpha = 1$, that is, slicing L along the vertical dotted line. The lower right panel shows the profile log likelihood for α , which corresponds to the log likelihood values along the dashed line in the panel above, plotted against α .



$$\theta = \hat{\theta}_\alpha$$

$$(\lambda = \hat{\lambda}_\alpha)$$

$$l_p(\hat{\alpha}) - l_p(\alpha)$$

$$l_p(\alpha)$$

$$\hat{\alpha} \pm 2(\hat{se})$$

$$-l_p^{cr}(\hat{\alpha})$$

$$l_p(\alpha) = l(\alpha, \hat{\lambda}_\alpha)$$

estimation

model $f(x; \theta)$ $\leftarrow$ density for the distⁿ of X , given θ

θ viewed as another r.v. $\left(\begin{array}{l} \textcircled{4} - \text{random} \\ \theta - \text{fixed} \end{array} \right)$

prior $\pi(\theta)$
density ~ prob ~ fr sp

posterior $\pi(\theta | x) = f(x | \theta) \pi(\theta) / \int f(x | \theta) \pi(\theta) d\theta$

sample $x_1, \dots, x_n$ note $f(\underline{x}; \theta) = \prod_{i=1}^n f(x_i | \theta)$ Bayes th.
cond'l prob.

$$P(B|A) = \frac{P(B, A)}{P(A)}$$

$$\pi(\theta | \underline{x}) = \frac{\prod_{i=1}^n f(x_i; \theta) \pi(\theta)}{\int \dots d\theta}$$

$0 \leq \pi(\theta | \underline{x}) \leq 1$

marginal for $\underline{x}$

Frequentist and Bayesian contrast

$$L(\theta; \underline{x}) = c(\underline{x}) f(\underline{x}; \theta)$$

$$L(\theta_1; \underline{x}) / L(\theta_2; \underline{x})$$

Frequentist:

- There is a fixed parameter (unknown) we are trying to learn
- Our methods are evaluated using probabilities based on $f(x; \theta)$

$$E_{\theta}\{l'(\theta; \underline{x})\} = 0$$

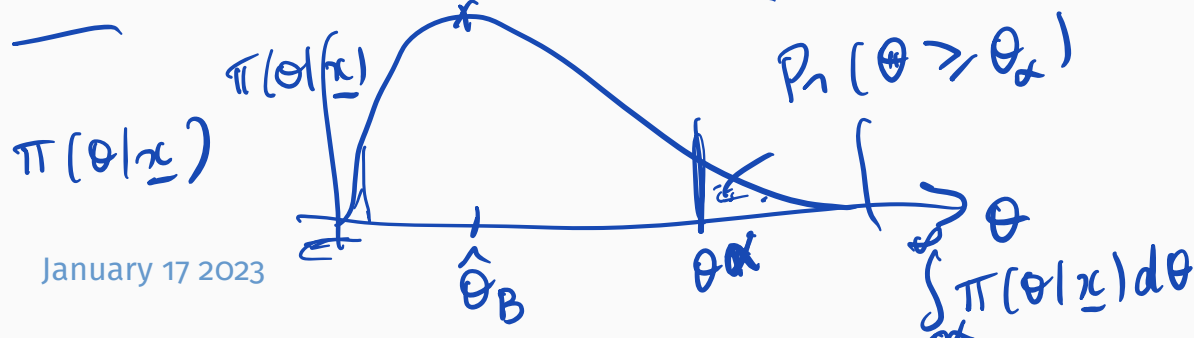
Bayesian:

- The parameter can be treated as a random variable ←
- We model its distribution $\pi(\theta)$ ← density
- Combine this with a model $f(\underline{x} | \theta)$ for data
- Update prior belief on the basis of the data ←

$$\int l'(\theta; \underline{x}) f(\underline{x}; \theta) d\underline{x}$$

$$l'(\hat{\theta}; \underline{x}) = 0$$

$$\hat{\theta} \xrightarrow{\uparrow \text{under } f(\underline{x}; \theta)} \theta$$



$X_1, \dots, X_n$ i.i.d. Bernoulli (θ)

$$\pi(\theta; \alpha, \beta) = \frac{\Gamma(\alpha + \beta)}{\Gamma(\alpha)\Gamma(\beta)} \theta^{\alpha-1} (1-\theta)^{\beta-1}, 0 < \theta < 1$$

posterior mean, mode

$$f(\underline{x}, \theta) = \prod_{i=1}^n \theta^{x_i} (1-\theta)^{1-x_i} = L(\theta, \underline{x}) = \theta^{\sum x_i} (1-\theta)^{n-\sum x_i}$$

$$0 \leq \theta \leq 1$$

$$\pi(\theta | \underline{x}) = \frac{\theta^{x_+} (1-\theta)^{n-x_+} \cdot \theta^{\alpha-1} (1-\theta)^{\beta-1}}{\int_0^1 \dots d\theta} \quad 0 \leq \theta \leq 1$$

$$= \theta^{x_+ + \alpha - 1} (1-\theta)^{n-x_+ + \beta - 1} / \int_0^1 \dots d\theta$$

$$\text{Be}(x_+ + \alpha, n - x_+ + \beta)$$

$$E(\text{Beta?}) = \frac{\alpha}{\alpha + \beta}$$

$$E \uparrow$$

$$(\sum x_i + \alpha) / (n + \alpha + \beta)$$

$$(\sum x_i) = \hat{\theta}$$

Example: censored exponential

$$\theta_B = E_{\pi(\theta|x)} \theta$$

MS 5.27

$X_1, \dots, X_n$ i.i.d. Exponential (λ)

censored at r smallest x ; let $Y_i = X_{(i)}, i = 1, \dots, r$

$$\pi(\lambda) \sim \text{Exp}(\alpha)$$

$$\alpha = \beta = 1$$

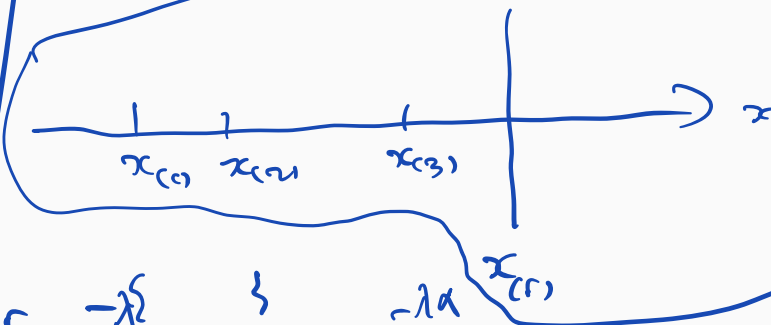
$$\hat{\theta}_B = \frac{\sum x_i + 1}{n + 2}$$

$$f(\mathbf{y} | \lambda) = \prod_{i=1}^r \lambda^r \exp(-\lambda y_i) \prod_{i=r+1}^n \exp(-\lambda y_i) = \lambda^r \exp\{-\lambda \sum_{i=1}^r y_i + (n-r)y_r\}$$

$$f(x_i; \lambda) = \lambda e^{-\lambda x_i}$$

$$1 - F(x; \lambda) = e^{-\lambda x}$$

$$\pi(\lambda | \mathbf{y}) = \lambda^r e^{-\lambda \sum_{i=1}^r y_i} \times \alpha e^{-\lambda x}$$



$$\pi(\theta | \underline{x}) = \text{Be}(x_+ + \alpha, n - x_+ + \beta)$$

$$E_{\pi(\theta | \underline{x})}(\theta) = \frac{x_+ + \alpha}{n + \beta + \alpha}$$

a possible est. of θ is

$$\hat{\theta}_B$$

or

$$\hat{\theta} = \frac{x_+}{n}$$

or mode of $\pi(\theta | \underline{x}) \equiv \text{argsup}_{\theta} \pi(\theta | \underline{x})$
 $\hat{\theta}_{B_2}$

$$\hat{\theta}_B = \left(\frac{x_+}{n} \right) w + \left(\frac{\alpha}{\alpha + \beta} \right) (1 - w)$$

$w = ??$ ntbc

↑
sample
average

└
prior exp'd value
of θ

$$\alpha = \beta = 1$$

$$\pi(\theta) = \begin{cases} 1 & 0 \leq \theta \leq 1 \\ 0 & \text{o.w.} \end{cases}$$

$$\frac{x_+ + 1}{n + 2}$$

$$f(x; \theta) = \exp\{c(\theta)T(x) - d(\theta) + S(x)\}; \quad \pi(\theta; \alpha, \beta) = K(\alpha, \beta) \exp\{\alpha c(\theta) - \beta d(\theta)\}$$

$$f(x; \theta) = \exp\{c(\theta)T(x) - d(\theta) + S(x)\}; \quad \pi(\theta; \alpha, \beta) = K(\alpha, \beta) \exp\{\alpha c(\theta) - \beta d(\theta)\}$$

Example: $f(x; \theta) = \theta(1 - \theta)^x, x = 0, 1, \dots; 0 < \theta < 1$

$$f(x; \theta) = \exp\{c(\theta)T(x) - d(\theta) + S(x)\}; \quad \pi(\theta; \alpha, \beta) = K(\alpha, \beta) \exp\{\alpha c(\theta) - \beta d(\theta)\}$$

Example: $f(x; \theta) = \theta(1 - \theta)^x, x = 0, 1, \dots; 0 < \theta < 1$

Example: $f(x; \mu) = \frac{1}{\sqrt{2\pi}} \exp\{-\frac{1}{2}(x - \mu)^2\}$

Table 3.1 Scores from two tests taken by 22 students, **mechanics** and **vectors**.

	1	2	3	4	5	6	7	8	9	10	11
mechanics	7	44	49	59	34	46	0	32	49	52	44
vectors	51	69	41	70	42	40	40	45	57	64	61

	12	13	14	15	16	17	18	19	20	21	22
mechanics	36	42	5	22	18	41	48	31	42	46	63
vectors	59	60	30	58	51	63	38	42	69	49	63

Table 3.1 shows the scores on two tests, **mechanics** and **vectors**, achieved by $n = 22$ students. The sample correlation coefficient between the two scores is $\hat{\theta} = 0.498$,

$$\hat{\theta} = \frac{\sum_{i=1}^{22} (m_i - \bar{m})(v_i - \bar{v})}{\left[\sum_{i=1}^{22} (m_i - \bar{m})^2 \sum_{i=1}^{22} (v_i - \bar{v})^2 \right]^{1/2}}, \quad (3.10)$$

with m and v short for **mechanics** and **vectors**, $\bar{m}$ and $\bar{v}$ their averages. We wish to assign a Bayesian measure of posterior accuracy to the true correlation coefficient θ , “true” meaning the correlation for the hypothetical population of all students, of which we observed only 22.

If we assume that the joint (m, v) distribution is bivariate normal (as

$$\begin{pmatrix} x_i \\ y_i \end{pmatrix} \sim N_2 \left(\begin{pmatrix} \mu_x \\ \mu_y \end{pmatrix}, \Sigma \right)$$

$$\theta = \text{corr}^t \text{ bet. } x \text{ \& } y$$

$$(\psi)$$

$$\frac{\sigma_{12}}{\sqrt{\sigma_1^2 \sigma_2^2}}$$

marginal for $\hat{\theta}$ in $N_2\left(\begin{pmatrix} \mu_1 \\ \mu_2 \end{pmatrix}, \Sigma\right)$

$$f(\hat{\theta} | \theta) = \frac{1}{\pi} (n-2)(1-\theta^2)^{(n-1)/2} (1-\hat{\theta}^2)^{(n-4)/2} \int_0^\infty \frac{1}{\cosh(w) - \theta\hat{\theta}} dw \quad ||$$

$\pi(\theta | \hat{\theta}) = \uparrow \times \pi(\theta) / \int \dots d\theta$

$\pi(\theta)$ prior $-1 \leq \theta \leq 1$

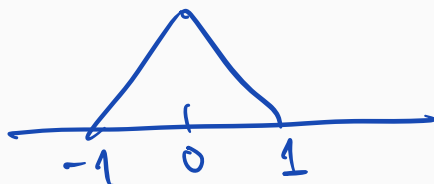
i) $\pi(\theta) = \frac{1}{2}$

← "ignorance"

ii) $\pi(\theta) = \frac{1}{1-\theta^2}, -1 \leq \theta \leq 1$

← Jeffreys'

iii)



←

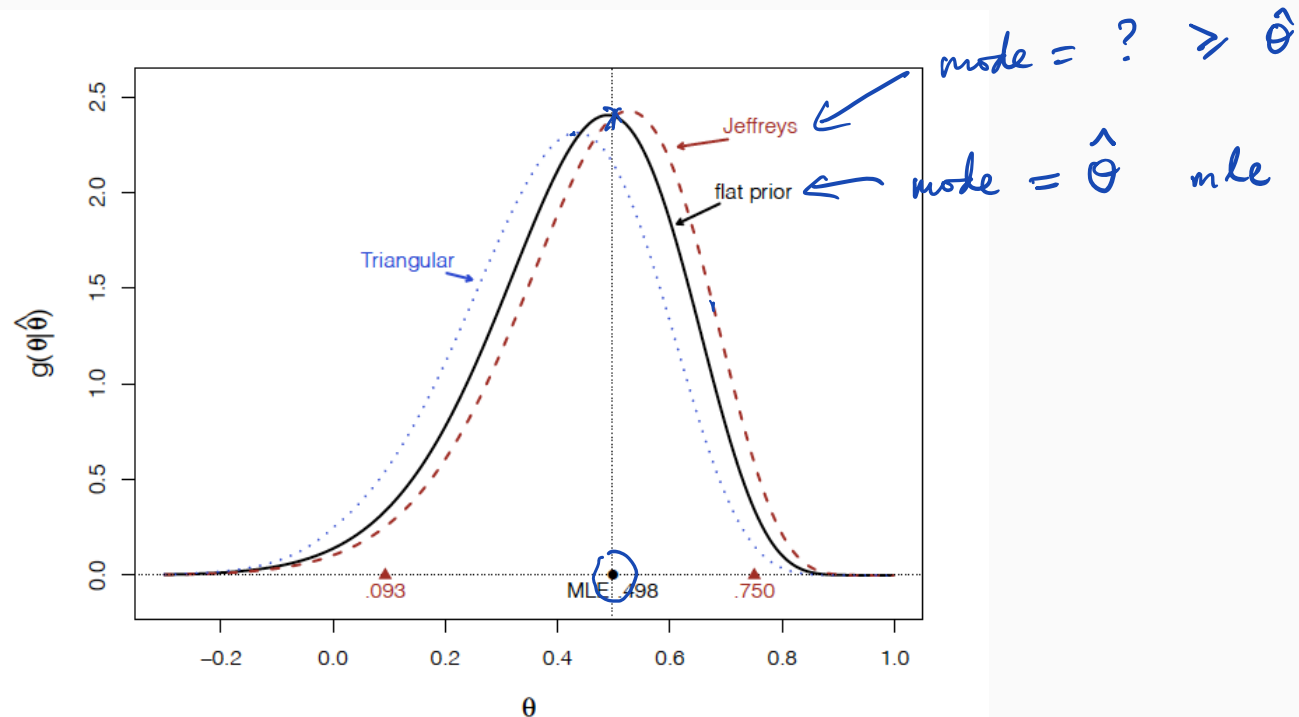


Figure 3.2 Student scores data; posterior density of correlation θ for three possible priors.

11.2 · Inference

579

Table 11.2 Mortality rates r/m from cardiac surgery in 12 hospitals (Spiegelhalter *et al.*, 1996b, p. 15). Shown are the numbers of deaths r out of m operations.

<i>A</i>	0/47	<i>B</i>	18/148	<i>C</i>	8/119	<i>D</i>	46/810	<i>E</i>	8/211	<i>F</i>	13/196
<i>G</i>	9/148	<i>H</i>	31/215	<i>I</i>	14/207	<i>J</i>	8/97	<i>K</i>	29/256	<i>L</i>	24/360

provided the mode lies inside the parameter space. Here $\tilde{J}(\theta)$ is the second derivative matrix of $\tilde{\ell}(\theta)$. This expansion corresponds to a posterior multivariate normal

prior for hospital A $Beta(1, 1)$

posterior mean

580

11 · Bayesian Models

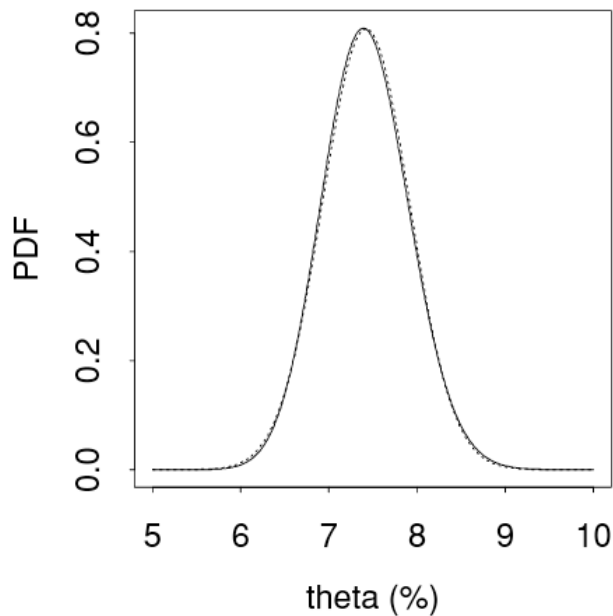
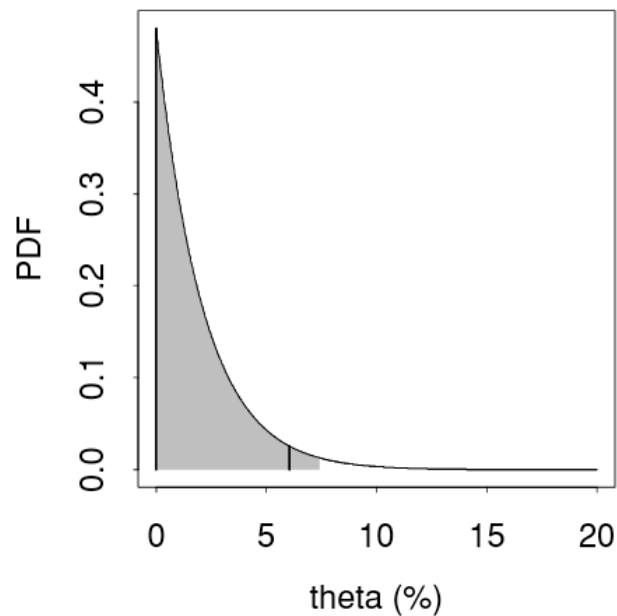


Figure 11.1 Cardiac surgery data. Left panel: posterior density for θ_A , showing boundaries of 0.95 highest posterior credible interval (vertical lines) and region between posterior 0.025 and 0.975 quantiles of $\pi(\theta_A | y)$ (shaded). Right panel: exact posterior beta density for overall mortality rate θ (solid) and normal approximation (dots).

put all hospitals together; 208 failures ‘

Marginalization

Not all likelihood functions are regular

Example: $X_1, \dots, X_n$ i.i.d. $U(o, \theta)$

... Not all likelihood functions are regular

MS Exercise 5.1

$X_1, \dots, X_n$ i.i.d. $f(x; \theta) = a(\theta_1, \theta_2)h(x), \quad \theta_1 \leq x \leq \theta_2$