

Mathematical Statistics II

STA2212H S LEC9101

Week 9

March 17 2021

Start recording!



Christopher Mims
@mims



This is one of the coolest and most creative visualizations of a phenomenon I've ever seen

It's also an important account of one of the biggest potential climate tipping points we may be rapidly pushing ourselves over

[nytimes.com/interactive/20...](https://www.nytimes.com/interactive/2021-03-04/climate/ice-sheet-collapse.html)

2021-03-04, 6:18 AM

link

- overview of Bayesian inference
- posterior predictive distribution ←
- Bayesian computation
- Example 11.9 AoS x10.8
- Bayesian hierarchical models

$$p(\underline{x}_{n+1} | \underline{x}^n)$$

$$= \int \underbrace{f(x_{n+1} | \theta)}_{\substack{\uparrow \\ \text{ass'g } x_{n+1} \perp\!\!\!\perp \underline{x}_n}} \pi(\theta | \underline{x}^n) d\theta$$

$$p(\underline{x}_{n+1}) = \int f(x_{n+1} | \theta) \pi(\theta) d\theta$$

↑
sometimes used for
planning studies

- overview of Bayesian inference
- posterior predictive distribution
- Bayesian computation
- **Example 11.9 AoS**
- Bayesian hierarchical models

Addendum: If $X_1, \dots, X_n$ are i.i.d. $N(\mu, \sigma^2)$, both parameters unknown, then the conjugate priors for μ and σ^2 are

$$\mu \sim N(\mu_0, \tau^2); \quad \sigma^2 \sim \text{IG}\left(\frac{\nu_0}{2}, \frac{\nu_0 \sigma_0^2}{2}\right)$$

Inverse Gamma distribution $\text{IG}(\alpha, \beta)$ has density $f(t; \alpha, \beta) = \frac{\beta^\alpha}{\Gamma(\alpha)} \frac{1}{x^{\alpha+1}} e^{-\beta/x}$

$$\propto x^{\alpha-1} e^{-\beta/x}$$

1. Notes from March 12 (Friday) still to be posted
2. Friday March 19 – Stein's paradox
3. empirical Bayes
4. introduction to decision theory

B-H proof

- Mar 22 3.00 – 4.00 pm EDT

[Data Science ARES](#)

Jesse Cisewski Kehe

University of Wisconsin-Madison

“Astrostatistics: From Exoplanets to
the Large-scale Structure of the Universe”



• $x_i | \theta_i \sim N(\theta_i, v_i)$, independent; $\theta_i | \mu \sim N(\mu, \sigma^2), i = 1, \dots, n$ i.i.d;

$\mu \sim N(\mu_0, \tau^2)$

$\hat{\theta}_i = x_i$

hyper-par

hyper-prior

$$E(\theta_i | \underline{x}^n) = x_i \left(\frac{\sigma^2}{\sigma^2 + v_i} \right) + \mu \left(\frac{1}{\sigma^2 + v_i} \right)$$

$v_i, \sigma^2, \mu_0, \tau^2$
all known

est. $\tilde{\theta}_i = x_i \left(\frac{\sigma^2}{\sigma^2 + v_i} \right) + \tilde{\mu} \left(\frac{1}{\sigma^2 + v_i} \right)$

$$\cancel{E(\mu | \underline{x}^n) = \frac{\sum x_i / (v_i + \sigma^2)}{\sum 1 / (v_i + \sigma^2)} = \tilde{\mu}}$$

shrinkage
estimators

If σ^2 unknown, no explicit solⁿ for $\tilde{\theta}_i$ & $\tilde{\mu}$: need use MCMC

- $x_i \mid \theta_i \sim N(\theta_i, v_i)$, independent; $\theta_i \mid \mu \sim N(\mu, \sigma^2), i = 1, \dots, n$ i.i.d; $\mu \sim N(\mu_0, \tau^2)$
- $v_i, i = 1, \dots, n, \sigma^2, \mu_0, \tau^2$ all known

- $x_i \mid \theta_i \sim N(\theta_i, v_i)$, independent; $\theta_i \mid \mu \sim N(\mu, \sigma^2)$, $i = 1, \dots, n$ i.i.d; $\mu \sim N(\mu_0, \tau^2)$
- $v_i, i = 1, \dots, n, \sigma^2, \mu_0, \tau^2$ all known

•

$$E(\theta_i \mid \mathbf{x}) = x_i \frac{\sigma^2}{\sigma^2 + v_i} + E(\mu \mid \mathbf{x}) \left(1 - \frac{\sigma^2}{\sigma^2 + v_i}\right)$$

- $x_i \mid \theta_i \sim N(\theta_i, v_i)$, independent; $\theta_i \mid \mu \sim N(\mu, \sigma^2)$, $i = 1, \dots, n$ i.i.d; $\mu \sim N(\mu_0, \tau^2)$
- $v_i, i = 1, \dots, n, \sigma^2, \mu_0, \tau^2$ all known

•

$$E(\theta_i \mid \mathbf{x}) = x_i \frac{\sigma^2}{\sigma^2 + v_i} + E(\mu \mid \mathbf{x}) \left(1 - \frac{\sigma^2}{\sigma^2 + v_i}\right)$$

•

$$E(\mu \mid \mathbf{x}) = \frac{\mu_0/\tau^2 + \sum x_i/(\sigma^2 + v_i)}{1/\tau^2 + \sum 1/(\sigma^2 + v_i)}$$

- $x_i \mid \theta_i \sim N(\theta_i, v_i)$, independent; $\theta_i \mid \mu \sim N(\mu, \sigma^2)$, $i = 1, \dots, n$ i.i.d; $\mu \sim N(\mu_0, \tau^2)$
- $v_i, i = 1, \dots, n, \sigma^2, \mu_0, \tau^2$ all known

•

$$E(\theta_i \mid \mathbf{x}) = x_i \frac{\sigma^2}{\sigma^2 + v_i} + E(\mu \mid \mathbf{x}) \left(1 - \frac{\sigma^2}{\sigma^2 + v_i}\right)$$

•

$$E(\mu \mid \mathbf{x}) = \frac{\mu_0/\tau^2 + \sum x_i/(\sigma^2 + v_i)}{1/\tau^2 + \sum 1/(\sigma^2 + v_i)}$$

- If σ^2 unknown, then need to sample from the posterior, no closed form available

- $x_i \mid \theta_i \sim N(\theta_i, v_i)$, independent; $\theta_i \mid \mu \sim N(\mu, \sigma^2), i = 1, \dots, n$ i.i.d; $\mu \sim N(\mu_0, \tau^2)$
- $v_i, i = 1, \dots, n, \sigma^2, \mu_0, \tau^2$ all known

$$E(\theta_i \mid \mathbf{x}) = x_i \frac{\sigma^2}{\sigma^2 + v_i} + E(\mu \mid \mathbf{x}) \left(1 - \frac{\sigma^2}{\sigma^2 + v_i}\right)$$

$$E(\mu \mid \mathbf{x}) = \frac{\mu_0/\tau^2 + \sum x_i/(\sigma^2 + v_i)}{1/\tau^2 + \sum 1/(\sigma^2 + v_i)}$$

- If σ^2 unknown, then need to sample from the posterior, no closed form available
- Figure 11.11 applies similar ideas, plus sampling from the posterior, in logistic regression

$\hat{\mu}$: avg hyp
 $L(\mu \mid \tilde{x}^n)$

$m(\tilde{x}^n)$

622

11 · Bayesian Models

each hosp. sep.

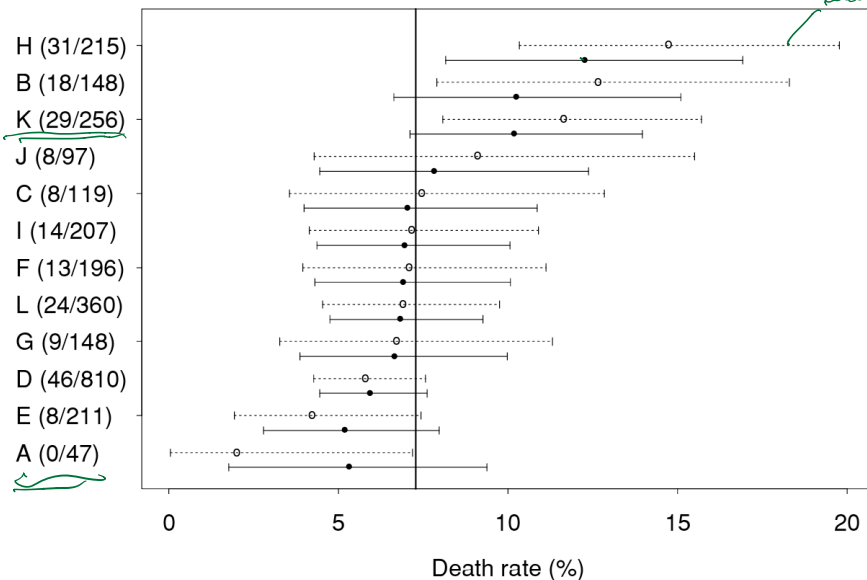


Figure 11.11 Posterior summaries for mortality rates for cardiac surgery data. Posterior means and 0.95 equitailed credible intervals for separate analyses for each hospital are shown by hollow circles and dotted lines, while blobs and solid lines show the corresponding quantities for a hierarchical model. Note the shrinkage of the estimates for the hierarchical model towards the overall posterior mean rate, shown as the solid vertical line; the hierarchical intervals are slightly shorter than those for the simpler model.

As above, $x_i | \theta_i \sim N(\theta_i, v_i)$, independent; $\theta_i | \mu \sim N(\mu, \sigma^2)$, $i = 1, \dots, n$ i.i.d;

$$\theta_i | x_i \sim N\left(x_i \frac{\sigma^2}{\sigma^2 + v_i} + \mu \frac{v_i}{\sigma^2 + v_i}, \frac{\sigma^2 v_i}{\sigma^2 + v_i}\right)$$

$$\pi(\theta_i | x_i) = \frac{f(x_i | \theta_i) \pi(\theta_i)}{\int f(x_i | \theta_i) \pi(\theta_i) d\theta_i} = \frac{m(x_i)}{f(x_i)}$$

marginal

$$\prod_{i=1}^n f_{x_i}(x_i | \mu) = L(\mu; \underline{x}^n)$$

(not bc ind't)

$$\hat{\mu} = \hat{\mu}(\underline{x}^n) = \frac{\sum x_i / (v_i + \sigma^2)}{\sum 1 / (v_i + \sigma^2)}$$

SM " x_i marginally ind't , $N(\mu, v_i / (\sigma^2 + v_i))$ "

$$\tilde{\theta}_i = \left(x_i \frac{\sigma^2}{\sigma^2 + v_i} \right) + \frac{\hat{\mu} \left(\frac{v_i}{\sigma^2 + v_i} \right)}{1} \\ = \hat{\mu} + (1 - \xi_i)(x_i - \hat{\mu})$$

$$\xi_i = \frac{v_i}{\sigma^2 + v_i} \quad \text{weight} \quad \begin{matrix} \sigma^2 \rightarrow \infty \\ v_i \rightarrow \infty \end{matrix}$$

Using ~~weight~~ ~~statistic~~ density for $\underline{x}^n$ to estimate parameters in the prior is called empirical Bayes

"large sets of parallel situations carry their own prior information"

$$\underline{E} \tilde{\theta}_i \neq \theta_i = \theta_i \frac{\sigma^2}{\sigma^2 + v_i} + \mu \frac{v_i}{\sigma^2 + v_i} \\ \text{bias} \neq 0 \quad \text{perhaps use } (\tilde{\theta}_i) \\ = E(\tilde{\theta}_i - \theta_i)^2 < \text{Var}(x_i)$$

$$x_i | \theta_i \sim N(\theta_i, v_i) \quad \Uparrow$$

Alternative to ridge regression

$$\min_{\beta} \left\{ \sum_{i=1}^n (x_i - z_i^T \beta)^2 + \lambda \sum_{j=1}^p \beta_j^2 \right\} \quad \begin{matrix} \beta_j = \text{change in } x_j \\ \text{per unit} \\ \text{change in } z_{ij} \end{matrix}$$

$\sum z_i = 0$
 $\sum z_i^2 = 1$ $\nearrow$
 $\leftarrow \frac{z_i - \bar{z}}{s_z}$

$$\hat{\beta}_{\text{bridge}} = (\underline{\underline{Z}}^T \underline{\underline{Z}} + \lambda \underline{\underline{I}})^{-1} \underline{\underline{Z}}^T \underline{\underline{y}} \quad \leftarrow \text{shrinkage est.'s}$$

76

Empirical Bayes

Table 6.1 *Counts y_x of number of claims x made in a single year by 9461 automobile insurance policy holders. Robbins' formula (6.7) estimates the number of claims expected in a succeeding year, for instance 0.168 for a customer in the $x = 0$ category. Parametric maximum likelihood analysis based on a gamma prior gives less noisy estimates.*

Claims x	0	1	2	3	4	5	6	7
Counts y_x	7840	1317	239	42	14	4	4	1
Formula (6.7)	.168	.363	.527	1.33	1.43	6.00	1.75	
Gamma MLE	.164	.398	.633	.87	1.10	1.34	1.57	

Handwritten annotations: Green arrows point to the first three rows. Green circles highlight the values .168, .164, .398, 6.00, 1.75, 1.34, and 1.57. A green arrow points to the 1.57 value. Green question marks are written below the 1.34 and 1.57 values.

$$f_k(x; \theta_k) = \frac{\theta_k^x e^{-\theta_k}}{x!} \quad x=0, 1, 2, \dots$$

k - customer k

$\nwarrow$ P_n " " had x claims in the year

$$\theta_k^0 e^{-\theta_k} / 0! = e^{-\theta_k}$$

$$\theta_k \sim \Gamma(\nu, 1/\sigma)$$

density $\theta_k^{\nu-1} e^{-\theta_k/\sigma} / \Gamma(\nu) \sigma^\nu$

marginal density $f(\underline{x}) = \prod_{k=1}^K \int f(\underline{x} | \theta_k) \pi(\theta_k) d\theta_k$

$$= f(\underline{x} | \nu, \sigma)$$

now use MLE to estimate τ, σ

$$L(\tau, \sigma | \underline{x}) \propto f(\underline{x} | \tau, \sigma)$$

estimates in the table are

$$E(\theta_i | \underline{x}) \quad \text{using } \hat{\sigma}, \hat{\tau} \text{ in the prior}$$

Other option is the prior is nonparametric

Nonparametric empirical Bayes

$$\rightarrow f(x) = \int_0^{\infty} \frac{e^{-\theta} \theta^x}{x!} \pi(\theta) d\theta$$

$$E(\theta|x) = \int \theta \underbrace{f(\theta|x) \pi(\theta) d\theta}_{f(x)} / \int \underbrace{f(\theta|x) \pi(\theta) d\theta}_{f(x)}$$

- Loss function $L(\theta; \hat{\theta}) : \Theta \times \hat{\Theta} \rightarrow \mathbb{R}$ θ unk. par. in model
 $\hat{\theta}$ estimate
- Risk of an estimator *reg.* $L(\theta; \hat{\theta}) = (\theta - \hat{\theta})^2$ sq. err
 $\mathbb{E} \quad |\theta - \hat{\theta}|$ abs. err
- Mean-squared error *class.* $\begin{cases} 0 & \hat{\theta} = \theta \\ 1 & \hat{\theta} \neq \theta \end{cases}$ (0-1 loss)
expected loss
- Bayes risk $|\theta - \hat{\theta}|^p$ L_p loss
- maximum risk Kullback-Leibler loss

"decision" is
 a choice of estimator

$$L(\theta, \hat{\theta}) = \int \left(\frac{f(x; \theta)}{f(x; \hat{\theta})} \right) f(x; \theta) dx$$

$$\hat{\theta} = \hat{\theta}(x).$$

- Loss function $L(\theta; \hat{\theta})$ $\swarrow$ unk. $\leftarrow$ choose
 - Risk of an estimator $R(\theta; \hat{\theta}) = \int L(\theta; \hat{\theta}) f(x; \theta) dx = E\{L(\theta, \hat{\theta}(x))\}$ $\uparrow$ est. $\theta \leftarrow$ undermodel
 - Mean-squared error $\uparrow$ under sq'd error loss
expected loss
 - Bayes risk
 - maximum risk
- $$\begin{aligned}
 R(\theta; \hat{\theta}) &= \int \{\theta - \hat{\theta}(x)\}^2 f(x; \theta) dx \\
 &= E\{\hat{\theta}(x) - \theta\}^2 = E_{\theta}\{\underbrace{\hat{\theta}(x) - E\hat{\theta}}_{+ E\hat{\theta} - \theta}\}^2 \\
 &= \text{var}\{\hat{\theta}(x)\} + \text{bias}^2\{\hat{\theta}(x)\}
 \end{aligned}$$

$$\bar{X} \sim \mathcal{N}(\theta, 1)$$

$$\hat{\theta}_1 = \bar{X}$$

$$\hat{\theta}_2 = 3$$

$$R(\theta; \hat{\theta}) = E\{\theta - \hat{\theta}(X)\}^2$$

$$1 = E(\theta - \bar{X})^2 = \frac{1}{n}$$

$$2 = E(\theta - 3)^2 = (\theta - 3)^2$$

12.2 Comparing Risk Functions

195

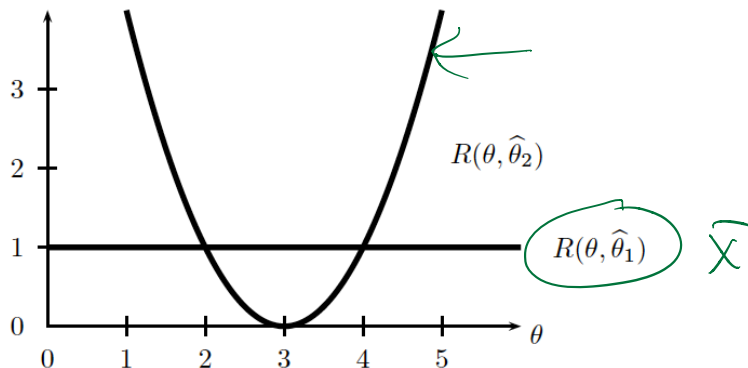


FIGURE 12.1. Comparing two risk functions. Neither risk function dominates the other at all values of θ .

often can't find an estimator
for which $R(\theta; \hat{\theta})$ is smallest $\forall \theta \in \Theta$

$$12.3 \quad X \sim \text{Bin}(n, p)$$

$$\hat{p}_1 = X/n = \bar{X}$$

$$R(p, \hat{p}_1) = \text{var} \bar{X} = \frac{p(1-p)}{n}$$

$$\hat{p}_2 = \frac{X + \alpha}{n + \alpha + \beta}$$

post. mean
if $B(\alpha, \beta)$

prior

$$R_{L_2}(p, \hat{p}_2) = \text{var} \hat{p}_2 + \{E(\hat{p}_2)\}^2$$

= see Eq. below Fig 12.2

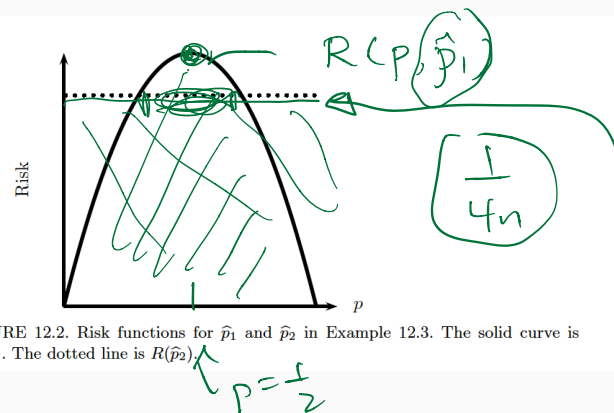


FIGURE 12.2. Risk functions for $\hat{p}_1$ and $\hat{p}_2$ in Example 12.3. The solid curve is $R(\hat{p}_1)$. The dotted line is $R(\hat{p}_2)$.

$$\alpha = \beta = \sqrt{n/4}$$

$$R_{L_2}(p, \hat{p}_2) = \frac{n}{4(n+\sqrt{n})^2}$$

$$R(\theta, \hat{\theta}(x)) = \int L(\theta, \hat{\theta}(x)) f(x; \theta) dx \quad \leftarrow \text{still dep}$$

$$R(\theta, d(x))$$

$$\hat{p}_1 = \bar{X};$$

$$R(p, \hat{p}_1) = \underline{p(1-p)/n}$$

$$\hat{p}_2 = \frac{n\bar{X} + \sqrt{n/4}}{n + \sqrt{n}};$$

$$R(p, \hat{p}_2) = \underline{\frac{n}{4(n + \sqrt{n})^2}}$$

Bin

$$\bar{R}(\hat{\theta}(x)) = \sup_{\theta \in \Theta} R(\theta, \hat{\theta}(x)) \quad \text{max risk}$$

minimax estimator is defined as $\inf_{\tilde{\theta} \in \tilde{\Theta}} \sup_{\theta \in \Theta} R(\theta, \tilde{\theta})$
 "minimax lower bd" "minimax decision rule"

Bayes rules

Bayes risk of $\hat{\theta}$

$$\int R(\theta, \hat{\theta}) \pi(\theta) d\theta = \mathbb{E}_{\pi}(\hat{\theta})$$
$$= r(f, \hat{\theta}) = r(\pi, \hat{\theta})$$

Bayes est. is the $f^* \hat{\theta}(x)$ that satisfies

$$\min_{\tilde{\theta} \in \mathcal{C}} r(f, \tilde{\theta}) = \min_{\tilde{\theta} \in \mathcal{C}} r_{\pi}(\tilde{\theta})$$

$$= \inf_{\tilde{\theta}} \mathbb{E}_{\pi}(\tilde{\theta})$$

Often called Bayes rules (decisions $\supset$ est's)

How to find Bayes rules? (i.e. Bayes estimators)

$$\text{def-: } \inf_{\tilde{\theta}} r_{\pi}(\tilde{\theta}) = \inf_{\tilde{\theta}} \int_{\Theta} R(\theta, \tilde{\theta}(x)) \pi(\theta) d\theta$$

Theorem 12.7 : note typo in stmt,
 $\theta \leftarrow \hat{\theta}$

Exp'd loss ^{under} posterior

$$= \int L(\theta, \hat{\theta}(x)) \pi(\theta|x) d\theta \quad \text{dep. } \hat{\theta}, x$$

Choose $\hat{\theta}_\pi(x)$ to $\min_{\hat{\theta}} \int L(\theta, \hat{\theta}(x)) \pi(\theta|x) d\theta$

What's the Bayes risk of $\hat{\theta}_\pi(x)$?

$$= \int R(\theta, \hat{\theta}_\pi(x)) \pi(\theta) d\theta = r_\pi(\hat{\theta}) \quad \text{dep. on } \pi$$

$$\min_{\tilde{\theta}} \int R(\theta, \tilde{\theta}_\pi(x)) \pi(\theta) d\theta \quad \leftarrow$$

$$= \min_{\tilde{\theta}} \left\{ \int L(\theta, \tilde{\theta}_\pi(x)) f(x|\theta) dx \right\} \pi(\theta) d\theta$$

$$f(x|\theta)\pi(\theta) = \pi(\theta|x) f(x)$$

$$= \min_{\tilde{\theta}} \int \int L(\theta, \tilde{\theta}_\pi(x)) \pi(\theta|x) f(x) dx d\theta$$

$$= \min_{\tilde{\theta}} \int f(x) \int L(\theta, \tilde{\theta}_\pi(x)) \pi(\theta|x) d\theta dx$$

$$\text{only need to } \min_{\tilde{\theta}} \int L(\theta, \tilde{\theta}_\pi(x)) \pi(\theta|x) d\theta \quad \leftarrow$$

Conclusion: "Bayes rule" re prior = best posterior estimate

$$\min_{\tilde{\theta}} \int (\theta - \tilde{\theta})^2 \pi(\theta|x) d\theta$$

$$= \dots = \hat{\theta} = E(\theta|x) \quad \text{answer}$$

admissibility $\leftarrow$ needed re Friday

2

11

2