

STA 2212S Mar 19, 2021

Recall Example 11.28 in SM Empirical Bayes

$$\Rightarrow x_i | \theta_i \sim N(\theta_i, v_i) \quad \text{ind't} \quad v_i \text{ known}$$

$$\Rightarrow \theta_i | \mu \sim N(\mu, \tau^2) \quad \text{i.i.d.} \quad \tau^2 \text{ known}$$

no hyperprior use data to est μ

$$\Rightarrow x_i \sim N(\mu, v_i + \tau^2) \quad \text{ind't} \quad E(x_i) = E\{E(x_i | \theta_i)\} = E(\theta_i) = \mu$$

$$L_m(\mu) \propto \prod_{i=1}^n f(x_i; \mu, v_i + \tau^2) \Rightarrow \mu \quad \text{var}(x_i) = E \text{var}(x_i | \theta_i) + \text{var}\{E(x_i | \theta_i)\}$$

$$\hat{\mu}_{EB} = \frac{\sum_{i=1}^n \{x_i / (v_i + \tau^2)\}}{\sum_{i=1}^n \{1 / (v_i + \tau^2)\}}$$

$$\hat{\theta}_{i,EB} = \hat{\mu}_{EB} + \frac{\tau^2}{v_i + \tau^2} (x_i - \hat{\mu}_{EB}) \quad \text{shrinkage est'r}$$

$$E(\theta_i | x_i) = \mu + \frac{\tau^2}{v_i + \tau^2} (x_i - \mu)$$

$$v_i \equiv 1$$

~~$$x_i \sim N(\mu, 1)$$~~

$$x_i | \theta_i \sim N(\theta_i, 1)$$

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

$$\Rightarrow x_i \sim N(\mu, \tau^2 + 1) \quad \text{var} \quad i=1, \dots, n \quad \text{i.i.d.}$$

n obs =
n par. unkn.
"nonparametric"

$$\hat{\mu}_{EB} = \bar{x}$$

$$\hat{\theta}_{i,EB} = \bar{x} + \frac{\tau^2}{\tau^2 + 1} (x_i - \bar{x})$$

Going further, τ^2 unknown $\Rightarrow \frac{1}{n} \sum (x_i - \bar{x})^2 = \frac{1}{n} \tau^2 + 1$

max. lik. est. of τ^2 [unbiased est. ...]

$$\hat{\tau}^2 = \left(\frac{1}{n} W - 1 \right)_+ \quad W = \sum (x_i - \bar{x})^2$$

$$\begin{aligned} \hat{\theta}_{i,EB} &= \bar{x} + \left(\frac{\frac{1}{n}W - 1}{\frac{1}{n}W} \right) (x_i - \bar{x}) \\ &= \bar{x} + \left\{ \frac{1 - (\frac{1}{n}W)^{-1}}{1} \right\} (x_i - \bar{x}) \\ &= \bar{x} + \left(1 - \frac{n}{W} \right) (x_i - \bar{x}) \end{aligned}$$

Instead, I'll use

$$\rightarrow \hat{\theta}_{i,EB} = \bar{x} + \left(1 - \frac{b}{W} \right) (x_i - \bar{x}) \quad b \text{ to be chosen}$$

b to minimize $E \sum_{i=1}^n (\hat{\theta}_{i,EB} - \theta_i)^2$

$R_b(\underline{\theta}, \hat{\underline{\theta}}) = \int L(\underline{\theta}, \hat{\underline{\theta}}(x)) f(x; \underline{\theta}) dx$ risk = exp'd loss

$$\min_b E \sum_{i=1}^n \left\{ \bar{x} + \left(1 - \frac{b}{W} \right) (x_i - \bar{x}) - \theta_i \right\}^2$$

$$\min_b E \sum_{i=1}^n \left\{ x_i - \theta_i - \frac{b}{W} (x_i - \bar{x}) \right\}^2$$

$$= \min_b \left\{ \sum \left\{ E(x_i - \theta_i)^2 - \sum \frac{2b}{W} (x_i - \theta_i)(x_i - \bar{x}) + \frac{b^2}{W^2} \sum (x_i - \bar{x})^2 \right\} \right\}$$

$$= \min_b \left\{ -E \sum \frac{2b}{W} (x_i - \theta_i)(x_i - \bar{x}) + E \frac{b^2}{W} \right\}$$

$$W = \sum (x_i - \bar{x})^2 \quad x_i \sim N(\mu, \sigma^2 + 1)$$

$$\sum (x_i - \bar{x})^2 \sim \sigma^2 \chi_{n-1}^2$$

$$\sim (\sigma^2 + 1) \chi_{n-1}^2$$

$$E\left(\frac{1}{W}\right) \quad E\left\{\frac{(x_i - \theta_i)(x_i - \bar{x})}{W}\right\} \quad \text{but harder}$$

" easy

$$= \min_b \left\{ n + b^2 E\left(\frac{1}{W}\right) - 2E\left[\sum_{i=1}^n \frac{(x_i - \bar{x})(x_i - \theta_i)}{W}\right] \right\}$$

$$= \dots \min_b \left\{ n + \underline{b(b - 2(n-3))} E(W^{-1}) \right\}$$

$$\frac{\partial}{\partial b}: \quad \{2b - 2(n-3)\} E(W^{-1}) = 0$$

$$b = n-3$$

$$\hat{\theta}_{i,JS} = \bar{x} + \left(1 - \frac{n-3}{W}\right)(x_i - \bar{x}) \quad n \geq 4$$

James - Stein estimator

for this argument

this has smaller $E(\text{mse})$ than $\underline{x_i}$

$$R(\underline{\theta}, \hat{\underline{\theta}}_{JS}) \leq R(\underline{\theta}, \underline{x})$$

~~in~~

if $X \sim N(\theta, 1)$ then X is best est.

$\mathbb{R}^1$

$x_1, \dots, x_n$ iid $N(\theta, 1)$ then $\bar{x}$ is best est.

$$\mathbb{R}^2 \quad \text{if } (x_{1i}, x_{2i})^T \sim N_2 \left(\begin{pmatrix} \theta_1 \\ \theta_2 \end{pmatrix}, \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \right) \quad i=1, \dots, n$$

then $(\bar{x}_1, \bar{x}_2)^T$ has smallest var

$$\mathbb{R}^3 \quad ??$$

in $\mathbb{R}^4$ $\underline{\bar{x}} = (\bar{x}_1, \dots, \bar{x}_4)$ does not have smallest var

HW6/7 bonus

$$X_i \sim N(\mu_i, 1)$$

$$\pi(\mu_i) \propto 1 \quad \text{"flat"}$$

$$\pi(\mu_i | x_i) \propto \cancel{N(x_i, 1)} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(\mu_i - x_i)^2}$$

$$\pi(\underline{\mu} | \underline{x}) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(\mu_i - x_i)^2}$$

But if you're interested in $\sum \mu_i^2$

$$E(\sum \mu_i^2 | \sum x_i^2) = n + \sum x_i^2 \quad \leftarrow$$

A frequentist would note

$$E(\sum x_i^2 | \underline{\mu}) = n + \sum \mu_i^2 \quad \leftarrow$$

$$\underline{\sum x_i^2 - n} \quad \text{is unbiased for } \sum \mu_i^2$$

Bayes posterior mean = $\sum x_i^2 + n = \tilde{\mu}_{\text{Bayes}}$

$$E(\tilde{\mu}_{\text{Bayes}} | \underline{x}) \text{ under model} = \sum \mu_i^2 + 2n$$

$$\hat{\theta}_{i,JS}^{\text{PP}} = \bar{x} + \left(1 - \frac{n-3}{\sum (x_i - \bar{x})^2}\right)_+ (x_i - \bar{x})$$

$$\hat{\theta}_{JS} = (\hat{\theta}_{1,JS}, \dots, \hat{\theta}_{n,JS}) \quad \text{smaller use than}$$

$$\hat{\theta}_{MLE} = (x_1, \dots, x_n)$$

As (12.12)

$$\hat{\theta}_{i,JS_2} = \left(1 - \frac{n-2}{\sum x_i^2}\right)_+ x_i \quad ??$$

$$x_i | \theta_i \sim N(\theta_i, 1) \quad \theta_i \sim N(0, \tau^2)$$

JS Then is $(\hat{\theta}_{1,JS_2}, \dots, \hat{\theta}_{n,JS_2})$ has smaller use than $(x_1, \dots, x_n)$

"positive part J-S estimator"