

Fisher scoring

Recap: Fisher information

Def: Let $\ell(\theta|Y)$ be a log likelihood,
and $u(\theta) = \frac{\partial \ell}{\partial \theta}$. The Fisher information

is

$$\hat{I}(\theta) = \text{Var}(u(\theta) | \theta)$$

$$\mathbb{E}[g(Y)] = \int g(y) f(y) dy$$

Properties

$$g(Y) = u(\theta)$$

Theorem: Under appropriate regularity conditions,

$$\mathbb{E}[U(\theta) | \theta] = 0$$

Proof: $\mathbb{E}[u(\theta) | \theta] = \mathbb{E}\left[\frac{\partial}{\partial \theta} \ell(\theta | Y) | \theta\right] = \mathbb{E}\left[\frac{\partial}{\partial \theta} \log f(Y | \theta) | \theta\right]$

$$= \mathbb{E}\left[\frac{1}{f(Y | \theta)} \cdot \frac{\partial}{\partial \theta} f(Y | \theta) | \theta\right] = \int \frac{\partial}{\partial \theta} f(Y | \theta) \frac{1}{f(Y | \theta)} f(Y | \theta) dy$$

$$= \int \frac{\partial}{\partial \theta} f(Y | \theta) dy$$

↘ required regularity:
swap derivative and integral
(see Casella & Berger 2.4)

$$= \frac{\partial}{\partial \theta} \int f(Y | \theta) dy$$

$$= \frac{\partial}{\partial \theta} 1 = 0 \quad //$$

$$\frac{\partial}{\partial \theta} \log f(\mathbf{y}|\theta) = \frac{\frac{\partial}{\partial \theta} f(\mathbf{y}|\theta)}{f(\mathbf{y}|\theta)}$$

Properties

Theorem: Under appropriate regularity conditions,

$$\mathcal{I}(\theta) = -\mathbb{E} \left[\frac{\partial^2}{\partial \theta^2} \ell(\theta|\mathbf{Y}) \mid \theta \right]$$

Proof:

$$\begin{aligned}
 & -\mathbb{E} \left[\frac{\partial^2}{\partial \theta^2} \ell(\theta|\mathbf{Y}) \mid \theta \right] = -\mathbb{E} \left[\frac{\partial^2}{\partial \theta^2} \log f(\mathbf{y}|\theta) \mid \theta \right] \\
 & = -\mathbb{E} \left[\frac{\partial}{\partial \theta} \left(\frac{\partial}{\partial \theta} \log f(\mathbf{y}|\theta) \right) \mid \theta \right] = -\mathbb{E} \left[\frac{\partial}{\partial \theta} \left(\frac{\frac{\partial}{\partial \theta} f(\mathbf{y}|\theta)}{f(\mathbf{y}|\theta)} \right) \mid \theta \right] \\
 & = -\mathbb{E} \left[\frac{\frac{\partial^2}{\partial \theta^2} f(\mathbf{y}|\theta)}{f(\mathbf{y}|\theta)} - \left(\frac{\frac{\partial}{\partial \theta} f(\mathbf{y}|\theta)}{f(\mathbf{y}|\theta)} \right)^2 \mid \theta \right] \\
 & = -\mathbb{E} \left[\underbrace{\frac{\frac{\partial^2}{\partial \theta^2} f(\mathbf{y}|\theta)}{f(\mathbf{y}|\theta)}}_{\frac{\partial^2}{\partial \theta^2} \int f(\mathbf{y}|\theta) d\mathbf{y} = 0} \mid \theta \right] + \mathbb{E} \left[\underbrace{\left(\frac{\partial}{\partial \theta} \log f(\mathbf{y}|\theta) \right)^2}_{\substack{= \mathbb{E}[(u(\theta))^2 \mid \theta] \\ = \text{Var}(u(\theta) \mid \theta) \\ = 0}} \mid \theta \right]
 \end{aligned}$$

Fisher information vs. Hessian

For logistic regression:

$$\hat{\mathcal{I}}(\beta) = X^T W X = -H(\beta)$$

under regularity conditions, $\hat{\mathcal{I}}(\beta) = -\mathbb{E}[H(\beta)]$

$\hat{\mathcal{I}}$ and $-H$ are not always the same

Example: $Y_1, \dots, Y_n \stackrel{\text{iid}}{\sim} \text{Bernoulli}(p)$

$$\frac{\partial^2}{\partial p^2} \ell(p|Y) = - \frac{\sum_i Y_i}{p^2} - \frac{(n - \sum_i Y_i)}{(1-p)^2} \quad \leftarrow \text{depends on the observed data}$$

$$\hat{\mathcal{I}}(p) = -\mathbb{E}\left[\frac{\partial^2}{\partial p^2} \ell(p|Y)\right] = \frac{n}{p(1-p)} \quad \leftarrow \text{does not depend on the observed data}$$

In general, $H(\theta)$ can be more work to calculate
 $\Rightarrow$ typically use $\hat{\mathcal{I}}(\theta)$ in place of $H(\theta)$

Newton's method:

$$\beta^{(r)} - \underbrace{H^{-1}(\beta^{(r)})}_{\text{replace w/ } \mathcal{I}^{-1}(\beta^{(r)})} u(\beta^{(r)})$$

Fisher scoring

1) Start with initial guess $\beta^{(0)}$

2) Update: $\beta^{(r+1)} = \beta^{(r)} + \mathcal{I}^{-1}(\beta^{(r)}) u(\beta^{(r)})$

3) Stop when $\beta^{(r+1)} \approx \beta^{(r)}$

IRLS for logistic regression

$$\begin{aligned}\beta^{(r+1)} &= \beta^{(r)} + (X^T W^{(r)} X)^{-1} X^T (Y - p^{(r)}) \\ &= \underbrace{(X^T W^{(r)} X)^{-1} X^T W^{(r)} X}_{\text{identity matrix}} \beta^{(r)} + (X^T W^{(r)} X)^{-1} X^T (Y - p^{(r)})\end{aligned}$$

$$= (X^T W^{(r)} X)^{-1} X^T W^{(r)} \left(X \beta^{(r)} + (W^{(r)})^{-1} (Y - p^{(r)}) \right)$$

$$\Rightarrow \beta^{(r+1)} = (X^T W^{(r)} X)^{-1} X^T W^{(r)} \begin{matrix} Z^{(r)} \\ Z^{(r)} \end{matrix} \quad \leftarrow \text{working responses at iteration } r$$

Linear regression: $\hat{\beta} = (X^T X)^{-1} X^T Y$

• actually doing weighted least squares with weights $W^{(r)}$, responses $Z^{(r)}$, explanatory variables X

Suppose $Y = X\beta + \varepsilon$

$$\underbrace{W^{\frac{1}{2}} Y}_{Y_w} = \underbrace{W^{\frac{1}{2}} X}_{X_w} \beta + \underbrace{W^{\frac{1}{2}} \varepsilon}_{\varepsilon_w}$$

$$Y_w = X_w \beta + \varepsilon_w$$

$$\begin{aligned} \hat{\beta} &= (X_w^T X_w)^{-1} X_w^T Y_w \\ &= (X^T W X)^{-1} X^T W Y \end{aligned}$$

$\varepsilon \sim N(0, W^{-1})$ (allow non-constant variance)

$W = \text{diag}(w_1, \dots, w_n)$ $W^{\frac{1}{2}} = \text{diag}(\sqrt{w_i})$

$\varepsilon_w \sim N(0, I)$

$\text{var}(W^{\frac{1}{2}} \varepsilon) = W^{\frac{1}{2}} W^{-1} W^{\frac{1}{2}} = I$

A preview of Fisher information properties