

Inferência para CS

Tópico 11 - Eficiência e Otimalidade do MLE Multivariado

Renato Martins Assunção

DCC - UFMG

MLE multivariado

- Vamos considerar o caso em que o parâmetro $\theta = (\theta_1, \dots, \theta_k) \in \mathbb{R}^k$ é um vetor com $k > 1$.
- O MLE $\hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_k)$ é também um vetor de dimensão k .
- O elemento $\hat{\theta}_i$ é um estimador de θ_i (a entrada i do vetor θ).
- O MLE $\hat{\theta}$ é obtido resolvendo-se o sistema de equações não-lineares que resulta de igualar o gradiente $\frac{\partial \ell}{\partial \theta}$ da log-verossimilhança $\ell(\theta)$ a zero:

$$\frac{\partial \ell}{\partial \theta} = \begin{bmatrix} \partial \ell / \partial \theta_1 \\ \partial \ell / \partial \theta_2 \\ \vdots \\ \partial \ell / \partial \theta_k \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{bmatrix}$$

- Queremos estudar o comportamento estatístico do MLE.
- Veremos que ele é um estimador ótimo de θ .

MLE multivariado

- O MLE $\hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_k)$ é um vetor ALEATÓRIO.
- Ele depende dos dados estocásticos e por isto varia de amostra para amostra.
- Precisamos estudar este comportamento aleatório do MLE.
- Como veremos, o comportamento é similar ao do caso unidimensional.
- Se n não é muito pequeno, de forma aproximada, teremos o seguinte:
 - O MLE $\hat{\theta}$ é aproximadamente não viciado para estimar o vetor θ . Isto é, $\mathbb{E}(\hat{\theta}) \approx \theta$.
 - Aproximadamente, a matriz de variância e covariância de $\hat{\theta}$ é, num certo sentido matemático, a “menor matriz” possível para um estimador não-viciado, a matriz de informação de Fisher $I(\theta)$.
 - O MLE $\hat{\theta}$ possui distribuição gaussiana multivariada aproximadamente.

Matriz de informação de Fisher

- A matriz de variância e covariância de $\hat{\theta}$ está associada à matriz de informação de Fisher $I(\theta)$.
- Definição: A matriz de informação de Fisher $I(\theta)$ é uma matriz $k \times k$ cujo elemento ij é dado por

$$[I(\theta)]_{ij} = \mathbb{E} \left(\frac{\partial \ell}{\partial \theta_i} \frac{\partial \ell}{\partial \theta_j} \right)$$

onde $\ell = \ell(\theta)$ é a log-verossimilhança de θ .

- Por exemplo, no caso em que $\theta = (\theta_1, \theta_2, \theta_3) \in \mathbb{R}^3$ teremos

$$I(\theta) = \begin{bmatrix} \mathbb{E} \left(\frac{\partial \ell}{\partial \theta_1} \right)^2 & \mathbb{E} \left(\frac{\partial \ell}{\partial \theta_1} \frac{\partial \ell}{\partial \theta_2} \right) & \mathbb{E} \left(\frac{\partial \ell}{\partial \theta_1} \frac{\partial \ell}{\partial \theta_3} \right) \\ \mathbb{E} \left(\frac{\partial \ell}{\partial \theta_2} \frac{\partial \ell}{\partial \theta_1} \right) & \mathbb{E} \left(\frac{\partial \ell}{\partial \theta_2} \right)^2 & \mathbb{E} \left(\frac{\partial \ell}{\partial \theta_2} \frac{\partial \ell}{\partial \theta_3} \right) \\ \mathbb{E} \left(\frac{\partial \ell}{\partial \theta_3} \frac{\partial \ell}{\partial \theta_1} \right) & \mathbb{E} \left(\frac{\partial \ell}{\partial \theta_3} \frac{\partial \ell}{\partial \theta_2} \right) & \mathbb{E} \left(\frac{\partial \ell}{\partial \theta_3} \right)^2 \end{bmatrix}$$

Exemplo de $l(\theta)$

- Suponha que $Y_1, Y_2, \dots, Y_n$ sejam i.i.d. $N(\mu, \sigma^2)$.
- Temos $\theta = (\mu, \sigma^2) \in \mathbb{R}^2$
- A log-verossimilhança de θ baseada numa amostra é dada por

$$\ell(\theta) = -\frac{n}{2} \log(2\pi) - \frac{n}{2} \log(\sigma^2) - \frac{1}{2\sigma^2} \sum_i (y_i - \mu)^2$$

- O vetor gradiente $k \times 1$ da log-verossimilhança $\ell(\theta)$ é dado por

$$\frac{\partial \ell}{\partial \theta} = \begin{bmatrix} \partial \ell / \partial \mu \\ \partial \ell / \partial \sigma^2 \end{bmatrix} = \begin{bmatrix} \sum_i (y_i - \mu) / \sigma^2 \\ -\frac{n}{2} \frac{1}{\sigma^2} + \frac{1}{2(\sigma^2)^2} \sum_i (y_i - \mu)^2 \end{bmatrix}$$

- A partir deste vetor, podemos obter a matriz de informação de Fisher tomando os seus produtos cruzados.

Exemplo de $I(\theta)$

- Temos

$$\begin{aligned}
 I(\theta) &= \mathbb{E} \begin{bmatrix} \left(\frac{\partial \ell}{\partial \mu} \right)^2 & \left(\frac{\partial \ell}{\partial \mu} \frac{\partial \ell}{\partial \sigma^2} \right) \\ \left(\frac{\partial \ell}{\partial \sigma^2} \frac{\partial \ell}{\partial \mu} \right) & \left(\frac{\partial \ell}{\partial \sigma^2} \right)^2 \end{bmatrix} \\
 &= \mathbb{E} \begin{bmatrix} \left(\sum_i \frac{Y_i - \mu}{\sigma^2} \right)^2 & \left(\sum_i \frac{Y_i - \mu}{\sigma^2} \right) \left(-\frac{n}{2} \frac{1}{\sigma^2} + \frac{1}{2} \frac{\sum_i (Y_i - \mu)^2}{(\sigma^2)^2} \right) \\ \left(-\frac{n}{2} \frac{1}{\sigma^2} + \frac{1}{2} \frac{\sum_i (Y_i - \mu)^2}{(\sigma^2)^2} \right) \left(\sum_i \frac{Y_i - \mu}{\sigma^2} \right) & \left(-\frac{n}{2} \frac{1}{\sigma^2} + \frac{1}{2} \frac{\sum_i (Y_i - \mu)^2}{(\sigma^2)^2} \right)^2 \end{bmatrix}
 \end{aligned}$$

- Um cálculo de probabilidade laborioso permite obter cada uma das esperanças presentes na matriz $I(\theta)$ resultando no seguinte:

$$I(\theta) = \begin{bmatrix} \frac{n}{\sigma^2} & 0 \\ 0 & \frac{n}{2\sigma^4} \end{bmatrix}$$

Uma fórmula alternativa para $I(\theta)$

- De forma similar ao caso unidimensional, podemos calcular a informação de Fisher de forma alternativa, usando a matriz Hessiana (a matriz de derivadas parciais de segunda ordem de $\ell(\theta)$):

$$[I(\theta)]_{ij} = -\mathbb{E} \left(\frac{\partial^2 \ell}{\partial \theta_i \partial \theta_j} \right)$$

- Por exemplo, no caso em que $\theta = (\theta_1, \theta_2, \theta_3) \in \mathbb{R}^3$, por esta fórmula alternativa, temos

$$I(\theta) = - \begin{bmatrix} \mathbb{E} \left(\frac{\partial^2 \ell}{\partial \theta_1^2} \right) & \mathbb{E} \left(\frac{\partial^2 \ell}{\partial \theta_1 \partial \theta_2} \right) & \mathbb{E} \left(\frac{\partial^2 \ell}{\partial \theta_1 \partial \theta_3} \right) \\ \mathbb{E} \left(\frac{\partial^2 \ell}{\partial \theta_2 \partial \theta_1} \right) & \mathbb{E} \left(\frac{\partial^2 \ell}{\partial \theta_2^2} \right) & \mathbb{E} \left(\frac{\partial^2 \ell}{\partial \theta_2 \partial \theta_3} \right) \\ \mathbb{E} \left(\frac{\partial^2 \ell}{\partial \theta_3 \partial \theta_1} \right) & \mathbb{E} \left(\frac{\partial^2 \ell}{\partial \theta_3 \partial \theta_2} \right) & \mathbb{E} \left(\frac{\partial^2 \ell}{\partial \theta_3^2} \right) \end{bmatrix}$$

- Note o sinal negativo na frente da matriz. Além disso, podemos passar o $\mathbb{E}$ para fora da matriz: esperança de matriz é a esperança de cada entrada da matriz.

Matriz de informação de Fisher

- Considerando o exemplo anterior de v.a.'s i.i.d. com distribuição $N(\mu, \sigma^2)$, a partir do vetor gradiente obtemos a matriz Hessiana abaixo:

$$\begin{bmatrix} -\frac{n}{\sigma^2} & -\frac{\sum_i (y_i - \mu)}{(\sigma^2)^2} \\ -\frac{\sum_i (y_i - \mu)}{(\sigma^2)^2} & \frac{n}{2} \frac{1}{(\sigma^2)^2} - \frac{\sum_i (y_i - \mu)^2}{(\sigma^2)^3} \end{bmatrix}$$

- Para obter $I(\theta)$, substituímos as instâncias $\mathbf{y}$ pelas v.a.'s $\mathbf{Y}$ e tomamos o negativo da esperança da matriz hessiana.
- Resulta que, neste exemplo, tomar a esperança das expressões da matriz Hessiana é muito simples. $I(\theta)$.
- O resultado é o mesmo de antes:

$$I(\theta) = \begin{bmatrix} \frac{n}{\sigma^2} & 0 \\ 0 & \frac{n}{2\sigma^4} \end{bmatrix}$$

Forma vetorial de $I(\theta)$

- A matriz de informação $I(\theta)$ pode ser definida termo a termo (como fizemos) ou de forma vetorial.
- A primeira fórmula para $I(\theta)$, que especifica o elemento ij da matriz como $[I(\theta)]_{ij} = \mathbb{E}(\partial\ell/\partial\theta_i \cdot \partial\ell/\partial\theta_j)$ pode ser escrita de forma vetorial como a esperança da matriz que resulta ao multiplicar duas vezes o vetor gradiente de $\ell(\theta)$ de forma transposta:

$$\begin{aligned}
 I(\theta) &= \mathbb{E} \left(\left[\frac{\partial \ell}{\partial \theta} \right] \cdot \left[\frac{\partial \ell}{\partial \theta} \right]' \right) \\
 &= \mathbb{E} \left(\begin{bmatrix} \frac{\partial \ell}{\partial \theta_1} \\ \frac{\partial \ell}{\partial \theta_2} \\ \vdots \\ \frac{\partial \ell}{\partial \theta_k} \end{bmatrix} \cdot \left[\frac{\partial \ell}{\partial \theta_1}, \frac{\partial \ell}{\partial \theta_2}, \dots, \frac{\partial \ell}{\partial \theta_k} \right] \right)
 \end{aligned}$$

Exemplo de $I(\theta)$

- Considerando o exemplo de n v.a.'s i.i.d. com distribuição $N(\mu, \sigma^2)$ e $\theta = (\mu, \sigma^2)$ temos

$$\begin{aligned}
 I(\theta) &= \mathbb{E} \left[\begin{array}{c} \partial \ell / \partial \mu \\ \partial \ell / \partial \sigma^2 \end{array} \right] \cdot \left[\partial \ell / \partial \mu, \partial \ell / \partial \sigma^2 \right] \\
 &= \mathbb{E} \left[\begin{array}{c} -\frac{n}{2} \frac{1}{\sigma^2} + \frac{\sum_i (y_i - \mu) / \sigma^2}{2(\sigma^2)^2} \sum_i (y_i - \mu)^2 \end{array} \right] \cdot \left[\sum_i (y_i - \mu) / \sigma^2, -\frac{n}{2} \frac{1}{\sigma^2} + \frac{1}{2(\sigma^2)^2} \sum_i (y_i - \mu)^2 \right]
 \end{aligned}$$

Outra forma vetorial de $I(\boldsymbol{\theta})$

- Além da forma vetorial dada por

$$I(\boldsymbol{\theta}) = \mathbb{E} \left(\left[\frac{\partial \ell}{\partial \boldsymbol{\theta}} \right] \cdot \left[\frac{\partial \ell}{\partial \boldsymbol{\theta}} \right]' \right)$$

podemos obter $I(\boldsymbol{\theta})$ usando uma forma vetorial para expressar a matriz hessiana de derivadas segundas.

- Especificamente, temos

$$I(\boldsymbol{\theta}) = -\mathbb{E} \left[\frac{\partial}{\partial \boldsymbol{\theta}} \frac{\partial \ell}{\partial \boldsymbol{\theta}} \right]$$

Desigualdade de Cramer-Rao multivariada

- Quando T é um estimador não viciado de $\theta \in \mathbb{R}$ vimos que $\text{MSE}(T) = \mathbb{V}(T) \geq 1/I(\theta)$.
- Temos uma versão multivariada deste resultado para $\theta = (\theta_1, \dots, \theta_k) \in \mathbb{R}^k$.
- Seja $\mathbf{T} = \mathbf{T}(\mathbf{Y}) = (T_1, \dots, T_k) \in \mathbb{R}^k$ um estimador k -dimensional não-viciado de θ (isto é, $\mathbb{E}(\mathbf{T}) = \theta$).
- Então, para todo i , temos

$$\text{MSE}(T_i) = \mathbb{V}(T_i) \geq [I^{-1}(\theta)]_{ii}$$

que é o elemento ii da INVERSA da matriz de informação de Fisher.

Desigualdade de Cramer-Rao multivariada (OPCIONAL)

- O resultado multivariado que apresentamos é apenas parcial.
- O resultado completo é um pouco mais sutil e útil em alguns problemas estatísticos de estimação. e testes de hipóteses em experimentos planejados.
- Considere qualquer combinação linear dos k elementos do estimador não viciado $\mathbf{T} = \mathbf{T}(\mathbf{Y})$:

$$W = c_1 T_1 + c_2 T_2 + \dots + c_k T_k = (c_1, \dots, c_k)' \cdot (T_1, \dots, T_k) = \mathbf{c}' \mathbf{T}$$

onde $\mathbf{c} = (c_1, c_2, \dots, c_n)$ são constantes conhecidas.

- É fácil mostrar que W é estimador não-viciado para $c_1 \theta_1 + c_2 \theta_2 + \dots + c_k \theta_k = \mathbf{c}' \boldsymbol{\theta}$.
- Então temos

$$\text{MSE}(W) = \mathbb{V}(W) = \mathbb{V}(\mathbf{c}' \mathbf{T}) \geq \mathbf{c}' \mathbf{I}(\boldsymbol{\theta}) \mathbf{c}$$

MLE: Caso Multivariado

- Suponha que o parâmetro $\theta = (\theta_1, \dots, \theta_k)$ é um vetor com k componentes
- O MLE de θ é um vetor aleatório $\hat{\theta}$ de dimensão k .
- O resultado anterior sobre o MLE é válido em multi-dimensões:

$$\hat{\theta} \approx N_k(\theta, I^{-1}(\theta))$$

- Se n não é pequeno, temos aproximadamente:
 - O MLE possui distribuição de probabilidade normal multivariada;
 - é aproximadamente não-viciado (isto é, $\mathbb{E}(\hat{\theta}) \approx \theta$);
 - A matriz de covariância do MLE é aproximadamente igual ao inverso de $I(\theta)$, que é um limite ótimo para estimadores não-viciados.

MLE em regressão linear

- No modelo de regressão linear múltipla, temos v.a.'s $Y_1, \dots, Y_n$ que são independentes mas não são i.i.d.: $Y_i \sim N(\mathbf{x}'_i \boldsymbol{\beta}, \sigma^2)$.
- A versão matricial do modelo de regressão linear é

$$\mathbf{Y} = \mathbf{X} \boldsymbol{\beta} + \boldsymbol{\varepsilon}$$

- onde $\mathbf{Y}$ é vetor $n \times 1$, a matriz de desenho $\mathbf{X}$ é de dimensão $n \times p$ com $p - 1$ variáveis independentes e a coluna de 1's.
- O vetor $\boldsymbol{\varepsilon} = (\varepsilon_1, \dots, \varepsilon_n)'$ é composto de v.a.'s i.i.d. $N(0, \sigma^2)$.
- O parâmetro $\boldsymbol{\theta}$ é $\boldsymbol{\theta} = (\boldsymbol{\beta}, \sigma^2) = (\beta_0, \dots, \beta_{p-1}, \sigma^2)$
- Queremos estimar $\boldsymbol{\beta}$ e σ^2 .

MLE de β coincide com mínimos quadrados

- O MLE $\hat{\beta}$ de $\beta = (\beta_0, \dots, \beta_p)$ coincide com o estimador de mínimos quadrados.
- Para ver isto, considere a log-verossimilhança $\ell(\theta)$.
- As v.a.'s Y_i são independentes com distribuição $N(\mathbf{x}'_i \beta, \sigma^2)$.
- Então

$$\begin{aligned}\ell(\theta) &= \log \left(\prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} \exp \left(-\frac{1}{2\sigma^2} (y_i - \mathbf{x}'_i \beta)^2 \right) \right) \\ &= -\frac{n}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} \sum_i (y_i - \mathbf{x}'_i \beta)^2\end{aligned}$$

- É fácil perceber que, para qualquer valor fixo de σ^2 , o vetor β que maximiza $\ell(\theta)$ será aquele que minimiza a soma de quadrados $\sum_i (y_i - \mathbf{x}'_i \beta)^2$.
- Assim, o MLE de β é o vetor que minimiza a distância entre o vetor $\mathbf{Y}$ e a combinação linear $\mathbf{X} \beta$:

MLE de β coincide com mínimos quadrados

- Já aprendemos que a solução deste último problema é a combinação que produz a projeção ortogonal de $\mathbf{Y}$ no espaço vetorial das combinações lineares das colunas de $\mathbf{X}$:

$$\hat{\beta} = (\hat{\beta}_0, \dots, \hat{\beta}_p)' = (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{Y}$$

- $\hat{\beta}$ é uma função do vetor de dados aleatórios $\mathbf{Y}$ e da matriz de regressores $\mathbf{X}$ (que é considerada uma matriz de constantes conhecida).
- Se escrevermos a matriz $k \times n$ dada por $(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t$ por A podemos ver que $\hat{\beta} = A \mathbf{Y}$.
- Assim, cada elemento de $\hat{\beta}$ é uma combinação linear dos elementos do vetor $\mathbf{Y}$.

Resíduos

- O modelo prediz ou estima o valor de Y_i usando $\hat{\beta}$.
- O vetor $n \times 1$ com os valores preditos pelo modelo para os valores realmente observados $\mathbf{Y}$ são dados por

$$\hat{\mathbf{Y}} = \mathbf{X} \hat{\beta} = \mathbf{X} (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{Y}$$

- A diferença entre $\mathbf{Y}$ e a predição $\hat{\mathbf{Y}}$ forma o vetor de resíduos ou vetor de erros de predição $\mathbf{r} = \mathbf{Y} - \hat{\mathbf{Y}}$.
- O i -ésimo elemento do vetor de resíduos é dado por

$$r_i = Y_i - \mathbf{x}_i' \hat{\beta}$$

- Gráficos e análises com os resíduos são uma excelente maneira de identificar falhas nas suposições do modelo de regressão linear

Resíduos

- Seja RSS a soma de quadrados dos resíduos (residual sum of squares)

$$RSS = ||\mathbf{Y} - \hat{\mathbf{Y}}||^2 = \sum_{i=1}^n (Y_i - \hat{Y}_i)^2 = \sum_{i=1}^n (Y_i - \mathbf{x}_i' \hat{\boldsymbol{\beta}})^2$$

- O MLE de σ^2 é dado pela média dos resíduos ao quadrado:

$$\hat{\sigma}^2 = \frac{RSS}{n}$$

- Para verificar isto, veja que tomando $\boldsymbol{\beta} = \hat{\boldsymbol{\beta}}$ teremos a log-verossimilhança $\ell(\boldsymbol{\theta})$ igual a

$$\ell(\hat{\boldsymbol{\beta}}, \sigma^2) = -\frac{n}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} RSS$$

que é maximizada se tomarmos $\sigma^2 = RSS/n$ (basta derivar e igualar a zero).

$\hat{\beta}$ é não-viciado

- Revisão de probabilidade: A é matriz de constantes e $\mathbf{Y}$ é vetor aleatório de dimensões compatíveis. Então o vetor aleatório $A\mathbf{Y}$ tem um vetor esperado igual a $\mathbb{E}(A\mathbf{Y}) = A\mathbb{E}(\mathbf{Y})$ e a matriz de covariância é $\mathbb{V}(A\mathbf{Y}) = A\mathbb{V}(\mathbf{Y})A^t$.
- O MLE $\hat{\beta}$ é não-viciado para β pois

$$\begin{aligned}\mathbb{E}(\hat{\beta}) &= \mathbb{E}((\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{Y}) \\ &= (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbb{E}(\mathbf{Y}) \\ &= (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbb{E}(\mathbf{X} \beta + \varepsilon) \\ &= (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t (\mathbf{X} \beta + \mathbb{E}(\varepsilon)) \\ &= \beta + \mathbf{0} = \beta\end{aligned}$$

Covariância de $\hat{\beta}$

- É possível calcular a matriz de covariância de $\hat{\beta}$.
- Chame $A = (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t$ e então $\hat{\beta} = A \mathbf{Y}$.
- Como as v.a.'s Y_i são independentes, sua matriz de covariância é $\sigma^2 I_n$, um múltiplo da matriz identidade de ordem n .
- Então, usando o resultado de probab, temos

$$\mathbb{V}(\hat{\beta}) = A(\sigma^2 I_n)A^t = \sigma^2 (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{X} (\mathbf{X}^t \mathbf{X})^{-1} = \sigma^2 (\mathbf{X}^t \mathbf{X})^{-1}$$

Propriedades de $\hat{\sigma}^2 = \frac{RSS}{n}$

- Quanto ao MLE de σ^2 no modelo de regressão linear, podemos mostrar que

$$\mathbb{E}(\hat{\sigma}^2) = \mathbb{E}\left(\frac{RSS}{n}\right) = \frac{n - k - 1}{n} \sigma^2$$

- Portanto o MLE $\hat{\sigma}^2$ é viciado para estimar σ^2 mas seu vício desaparece quando n cresce.
- É comum usarmos um estimador não-viciado para σ^2 dados por

$$S^2 = \frac{RSS}{n - k - 1}$$

- O uso deste estimador não-viciado (que não é o MLE) permite fazer vários cálculos exatos de probabilidade envolvidos em intervalos de confiança e testes de hipótese (mais a frente no curso).
- Para o MLE $\hat{\sigma}^2 = RSS/n$ temos a variância $2\sigma^4(n - k - 1)/n^2$.
- Além disso, a covariância entre $\hat{\beta}$ e $\hat{\sigma}^2$ é o vetor $\mathbf{0}$.

$I(\theta)$

- Podemos voltar a log-verossimilhança $\ell(\theta)$ e calcular a matriz de informação de Fisher usando vários truques de probabilidade e álgebra linear.
- No final encontramos

$$I(\theta) = \begin{pmatrix} \frac{1}{\sigma^2} \mathbf{X}^t \mathbf{X} & \mathbf{0} \\ \mathbf{0}^t & \frac{n}{2\sigma^4} \end{pmatrix}$$

- A matriz inversa é igual a

$$I^{-1}(\theta) = \begin{pmatrix} \sigma^2 (\mathbf{X}^t \mathbf{X})^{-1} & \mathbf{0} \\ \mathbf{0}^t & \frac{2\sigma^4}{n} \end{pmatrix}$$

$I(\theta)$

- A cota-inferior de Cramér-Rao para um estimador não-viciado de θ é

$$I^{-1}(\theta) = \begin{pmatrix} \sigma^2(\mathbf{X}^t \mathbf{X})^{-1} & \mathbf{0} \\ \mathbf{0}^t & \frac{2\sigma^4}{n} \end{pmatrix}$$

- Compare $I^{-1}(\theta)$ com a matriz de covariância exata do MLE

$$\mathbb{V}(\hat{\theta}) = \mathbb{V}(\hat{\beta}, \hat{\sigma}^2) = \begin{pmatrix} \sigma^2(\mathbf{X}^t \mathbf{X})^{-1} & \mathbf{0} \\ \mathbf{0}^t & \frac{2\sigma^4(n-k-1)}{n^2} \end{pmatrix} \rightarrow I^{-1}(\theta)$$

- O MLE do parâmetro β é não-viciado para β e atinge a cota inferior de Cramér-Rao $\sigma^2(\mathbf{X}^t \mathbf{X})^{-1}$
- O MLE de σ^2 tem um vício que vai a zero com n e tem uma variância que aproxima-se rapidamente da cota inferior de Cramér-Rao.

MLE de regressão logística

- No caso da regressão logística, o MLE $\hat{\theta}$ vai satisfazer a equação de verossimilhança não-linear

$$\frac{\partial \ell(\theta)}{\partial \theta} = \mathbf{X}^t \mathbf{y} - \mathbf{X}^t \mathbf{p} = \mathbf{0}$$

onde $\mathbf{p} = (p_1, \dots, p_n)$ e $p_i = 1/(1 + \exp(-\mathbf{x}_i^t \theta))$.

- O MLE é o resultado da convergência do algoritmo de Newton-Raphson:

$$\theta^{\text{new}} = \theta^{\text{old}} - \left(\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^t} \right)^{-1} \frac{\partial \ell(\theta)}{\partial \theta}$$

onde as derivadas são avaliadas usando-se o valor corrente θ^{old} .

- Como valor inicial, podemos usar $\theta^{(0)} = \mathbf{0}$.

MLE de regressão logística

- Temos

$$\frac{\partial \ell(\boldsymbol{\theta})}{\partial \boldsymbol{\theta}} = \mathbf{X}^t \mathbf{y} - \mathbf{X}^t \mathbf{p}$$

- e

$$\frac{\partial^2 \ell(\boldsymbol{\theta})}{\partial \boldsymbol{\theta} \partial \boldsymbol{\theta}^t} = -\mathbf{X}^t \mathbf{W} \mathbf{X}$$

onde

$$\mathbf{W} = \begin{bmatrix} p_1(1-p_1) & 0 & 0 & \dots & 0 \\ 0 & p_2(1-p_2) & 0 & \dots & 0 \\ 0 & 0 & p_3(1-p_3) & \dots & 0 \\ \vdots & \vdots & \vdots & \dots & \vdots \\ 0 & 0 & 0 & \dots & p_n(1-p_n) \end{bmatrix}$$

- Temos então

$$\boldsymbol{\theta}^{\text{new}} = \boldsymbol{\theta}^{\text{old}} + (\mathbf{X}^t \mathbf{W} \mathbf{X})^{-1} \mathbf{X}^t (\mathbf{y} - \mathbf{p})$$

$I(\theta)$

- Para obter a matriz de informação de Fisher, calculamos a matriz Hessiana (a matriz de derivadas parciais de segunda ordem da log-verossimilhança $\ell(\theta)$):

$$\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^t}$$

e depois calculamos $I(\theta)$ como o negativo do valor esperado de cada elemento desta matriz hessiana.

- A matriz Hessiana é igual a

$$\frac{\partial^2 \ell(\theta)}{\partial \theta \partial \theta^t} = -\mathbf{X}^t \mathbf{W} \mathbf{X}$$

- Que sorte!! Esta é uma matriz não-aleatória, somente com constantes.
- Assim, sua esperança é igual a ela mesma!
- Portanto, a matriz de informação de Fisher é

$$I(\theta) = \mathbf{X}^t \mathbf{W} \mathbf{X}$$

Exemplo - EMV de logística

- Conclusão:

$$\hat{\theta} \sim N_{k+1}(\theta, (\mathbf{X}^t \mathbf{W} \mathbf{X})^{-1})$$

com

$$\mathbf{W} = \text{diag}(p_1(1 - p_1), p_2(1 - p_2), \dots, p_n(1 - p_n))$$

- Cada p_i é função das covariáveis e de θ :

$$p_i = 1/(1 + \exp(-\mathbf{x}_i^t \theta))$$

- Isto é, o EMV de θ tem uma distribuição aproximadamente normal multivariada de dimensão $k + 1$ centrada no verdadeiro valor do parâmetro θ e com matriz de covariância $I^{-1}(\theta) = (\mathbf{X}^t \mathbf{W} \mathbf{X})^{-1}$
- A matrix $\mathbf{W}$ depende do valor desconhecido do parâmetro θ .
- Nós podemos usar nosso (melhor) estimador disponível $\hat{\theta}$ em $\mathbf{W}$ para obter uma aproximação para $I(\theta)$.

Intervalos de Confiança com o MLE

- O MLE $\hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_k)$ é um vetor ALEATÓRIO.
- A sua distribuição conjunta é dada por

$$\hat{\theta} \sim N_k(\theta, \mathbf{I}^{-1}(\theta))$$

- Considerando apenas a i -ésima coordenada, temos

$$\hat{\theta}_i \approx N_1(\theta_i, [\mathbf{I}^{-1}(\theta)]_{ii})$$

onde $[\mathbf{I}^{-1}(\theta)]_{ii}$ é o elemento (i, i) na diagonal da INVERSA da matriz de informação de Fisher.

- Podemos fornecer intervalos de confiança para cada θ_i .

Intervalos de Confiança com o MLE

- Considerando apenas a i -ésima coordenada, temos

$$\hat{\theta}_i \approx N_1(\theta_i, [\mathbf{I}^{-1}(\boldsymbol{\theta})]_{ii})$$

- Numa gaussiana qualquer, a probabilidade da variável afastar-se de sua esperança por menos que dois desvios-padrão é aproximadamente 0.95.
- Portanto, aproximadamente,

$$\mathbb{P}\left(|\hat{\theta}_i - \theta_i| \leq 2\sqrt{[\mathbf{I}^{-1}(\boldsymbol{\theta})]_{ii}}\right) \approx 0.95$$

- Desigualdade básica: $|a - b| \leq 2$ se, e somente se, $a - 2 \leq b \leq a + 2$
- Assim,

$$\mathbb{P}\left(\hat{\theta}_i - 2\sqrt{[\mathbf{I}^{-1}(\boldsymbol{\theta})]_{ii}} \leq \theta_i \leq \hat{\theta}_i + 2\sqrt{[\mathbf{I}^{-1}(\boldsymbol{\theta})]_{ii}}\right) \approx 0.95$$

Intervalos de Confiança com o MLE

- O intervalo

$$\left[\hat{\theta}_i - 2\sqrt{[\mathbf{I}^{-1}(\boldsymbol{\theta})]_{ii}}, \leq \hat{\theta}_i + 2\sqrt{[\mathbf{I}^{-1}(\boldsymbol{\theta})]_{ii}} \right]$$

é chamado intervalo de confiança de 95% para o parâmetro θ_i .

- A confiança é no método. 95% dos intervalos que você fizer usando o MLE vão cobrir o verdadeiro valor do parâmetro.
- Em geral, $\mathbf{I}^{-1}(\boldsymbol{\theta})$ depende do parâmetro $\boldsymbol{\theta}$, que é desconhecido.
- Nestes casos, substitua $\boldsymbol{\theta}$ por $\hat{\boldsymbol{\theta}}$.