

Modèle linéaire et moindres carrés

I Cadre et hypothèses

- Variable d'intérêt $y \in \mathbb{R}$ (variable endogène)
Variable explicative $x \in \mathbb{R}^p$ (variable exogène) donnée

- But: prédire y en fonction de x
+ précisément on cherche $E(y|x)$

- Modèle linéaire $E(y|x) = x^* \beta$ avec $\beta_0 \in \mathbb{R}^p$

Soit en posant $\varepsilon = y - E(y|x)$

$$y = x^* \beta + \varepsilon \quad \text{avec} \quad E(\varepsilon|x) = 0$$

Hypothèse forte de structure!

- Pb: inférence sur β_0 à partir de $(y_i = x_i^* \beta + \varepsilon_i)_{1 \leq i \leq n}$

- Hypothèse H_ε sur les résidus

$$H_\varepsilon = \begin{cases} E(\varepsilon_i | x_i) = 0 & (1) \end{cases}$$

$$\begin{cases} \text{Var}(\varepsilon_i) = \sigma_0^2 & (2) \end{cases}$$

$$\begin{cases} \text{Cov}(\varepsilon_i, \varepsilon_j) = 0 \text{ si } i \neq j & (3) \end{cases}$$

(1) résidus centrés conditionnellement à x (indispensable)

(2) variance finie + (3) décorrélation (ε bruit blanc)

- $H_\varepsilon \Rightarrow$
 $E(y_i) = x_i^* \beta$
 $\text{Var}(y_i) = \sigma^2$
 $\text{Cov}(y_i, y_j) = 0$ si $i \neq j$
 $\triangleq$ conditionnellement à x_i implicite.

Pb inférence sur β et σ_0^2 (les paramètres du modèle)

- Notation matricielle

$$Y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} \quad X = \begin{pmatrix} x_1^* \\ \vdots \\ x_n^* \end{pmatrix} \quad \varepsilon = \begin{pmatrix} \varepsilon_1 \\ \vdots \\ \varepsilon_n \end{pmatrix}$$

Observation $Y = X\beta + \varepsilon$

$H_\varepsilon \Rightarrow E(\varepsilon) = 0 \quad \text{Cov}(\varepsilon) = \sigma^2 I_m$

$E(Y) = X\beta \quad \text{Cov}(Y) = \sigma^2 I_m$

Rq: Si on note $X = \begin{pmatrix} x_1^v \\ \vdots \\ x_m^v \end{pmatrix} = (x_1 \dots x_p)$

$X\beta = \sum x_e \beta_e = \sum \beta_e x_e$

■ Identifiabilité si il existe une solution unique à $X\beta = X\beta_0$.

Si l'on veut un résultat uniforme en β .

$\Rightarrow X$ est de rang $\geq p$ (le maximum possible $X \in \mathbb{R}_{m \times p}$)

Rq: X de rang $p \Rightarrow \text{Im } X$ est un sev de $\mathbb{R}^m$ de dimension p .

• Si pas d'identifiabilité, pb nul rés externe de β

$\Rightarrow$ chgt de paramétrisation (ajout de variables, regroupement, ...)

■ Matrice X a un caractère arbitraire: ordre des variables / ordre des données exemples

$Y = X\beta + \varepsilon \quad \text{et} \quad Y = \tilde{X}\tilde{\beta} + \varepsilon$

mt 2 paramétrisations équivalentes si l'on peut passer de β_0 à $\tilde{\beta}_0$ de manière unique.

Ex $\tilde{X} = XZ^{-1}$ et $\tilde{\beta} = Z\beta$ avec Z inversible

Même prédiction mais interprétation et sélection peuvent être $\neq$

Ex: régression affine simple

$y_i = \beta_0 + \beta_1 x_i + \varepsilon_i$

$X = \begin{pmatrix} 1 & x_1 \\ \vdots & \vdots \\ 1 & x_m \end{pmatrix} \quad \beta = \begin{pmatrix} \beta_0 \\ \beta_1 \end{pmatrix}$

Pourcent équivalente $y_i = \beta'_0 + \beta_1 (x_i - \bar{x}) + \varepsilon_i$

$(\beta'_0 - \beta_1 \bar{x} = \beta_0)$

II Moindre carré (ordinaire) MCO

(least square) sur X de rang plein (identifiable)

Idée : Chercher $\tilde{\beta}$ qui minimise la somme des carrés résiduels
(SCR en français SSR en anglais)

$$SCR(\tilde{\beta}) = \|Y - X\tilde{\beta}\|^2 = \sum (y_i - x_i^* \tilde{\beta})^2$$

Termes d'attachement aux données qui pénalise des prédictions proches des valeurs observées

$$P_4: E(SCR(\beta)) = E(\| \epsilon \|^2) = \text{Tr}(Cov(\epsilon)) = \sigma^2 n$$

on avait pu remplacer le norme 2 par une autre norme (ou une autre fonction) ($\| \cdot \|_1$ ou même $\| \cdot \|_1$ et $\| \cdot \|_2$ à la Huber)
mais calcul facile et bonne propriété théorique

Propriété 1) $\nabla SCR(\tilde{\beta}) = -2X^*(Y - X\tilde{\beta})$ (et $\nabla^2 SCR(\tilde{\beta}) = 2X^*X$)

2) $\frac{\partial^2 SCR}{\partial \tilde{\beta}^2}(\tilde{\beta}) = 2X^*X$

3) SCR est strictement convexe.

Preuve:

$$SCR(\tilde{\beta}) = \|Y - X\tilde{\beta}\|^2$$

$$\begin{aligned} SCR(\tilde{\beta} + h) &= \| (Y - X\tilde{\beta}) - Xh \|^2 \\ &= \|Y - X\tilde{\beta}\|^2 - 2 \langle Xh, Y - X\tilde{\beta} \rangle + \|Xh\|^2 \\ &= \|Y - X\tilde{\beta}\|^2 - 2 \langle h, X^*(Y - X\tilde{\beta}) \rangle + \|Xh\|^2 \end{aligned}$$

$$\Rightarrow \nabla SCR(\tilde{\beta}) = -2X^*(Y - X\tilde{\beta})$$

$$\Rightarrow \frac{\partial^2 SCR}{\partial \tilde{\beta}^2}(\tilde{\beta}) = 2X^*X$$

$$\begin{aligned} \text{or } \forall \alpha \in \mathbb{R}^p \quad \alpha^* (2X^*X) \alpha \\ = 2 (X\alpha)^* X\alpha = 2 \|X\alpha\|_2^2 \geq 0 \end{aligned}$$

et même > 0 car X de rang plein
 $\alpha \neq 0$

$\Rightarrow 2X^*X$ défini positif $\Rightarrow$ SCR strictement convexe

Cor. $\hat{\beta} = \underset{\tilde{\beta}}{\operatorname{argmin}} \operatorname{SCR}(\tilde{\beta}) = \underset{\tilde{\beta}}{\operatorname{argmin}} \|Y - X\tilde{\beta}\|^2$
 existe et est unique.

De plus $\hat{\beta} = (X^*X)^{-1} X^*Y$

Preuve. le premier point découle de la stricte convexité
 le second s'obtient en utilisant la condition d'ordre 1 du min

$$\nabla \operatorname{SCR}(\hat{\beta}) = 0 = -2X^*(Y - X\hat{\beta})$$

$$\Rightarrow X^*X\hat{\beta} = X^*Y$$

or X^*X définie positive $\Rightarrow X^*X$ inversible

$$\Rightarrow \hat{\beta} = (X^*X)^{-1} X^*Y$$

Proposition Estimation de β par moindres carrés

L'estimateur des moindres carrés $\hat{\beta} = \underset{\tilde{\beta}}{\operatorname{argmin}} \operatorname{SCR}(\tilde{\beta})$ est donné

par $\hat{\beta} = (X^*X)^{-1} X^*Y$

et vérifie 1) $E(\hat{\beta}) = \beta$

2) $\operatorname{Cov}(\hat{\beta}) = \sigma^2 (X^*X)^{-1}$

Preuve. la première partie est une résulte du th précédent

$$E(\hat{\beta}) = E((X^*X)^{-1} X^*Y)$$

$$= (X^*X)^{-1} X^* E(Y) = (X^*X)^{-1} X^*X\beta$$

$$= \beta$$

$$\operatorname{Cov}(\hat{\beta}) = \operatorname{Var}((X^*X)^{-1} X^*Y)$$

Pq si $\tilde{\beta} = Z\beta$ alors

par les moindres carrés

$$\hat{\tilde{\beta}} = Z\hat{\beta}$$

$$= (X^*X)^{-1} X^* \operatorname{Cov}(Y) ((X^*X)^{-1} X^*)^*$$

$$= (X^*X)^{-1} X^* (\sigma^2 I_n) X (X^*X)^{-1}$$

$$= \sigma^2 (X^*X)^{-1}$$

Prop (Gauss-Markov) Parmi les estimateurs linéaires sans biais de β ,
 l'estimateur des moindres carrés $\hat{\beta}$ est celui de moindre variance
 (BLUE Best Unbiased Linear Estimator)

Preuve : $\beta^* = LY$

$$E(\beta^*) = \beta \Rightarrow LX\beta = \beta \quad \forall \beta$$

$$\Rightarrow LX = Id$$

$$\text{Cov}(\beta^*) = L \text{Cov}(Y) L^* = \sigma^2 LL^*$$

$$= \sigma^2 (M + \Delta) (M + \Delta)^* \quad \text{où } M = (X^*X)^{-1} X^*$$

$$\text{et } \Delta = L - M$$

$$= \sigma^2 [M M^* + \Delta \Delta^* + M \Delta^* + \Delta M^*]$$

or $\Delta M^* = \underbrace{(L - M) X (X^*X)^{-1}}_{= 0 \text{ car } L \text{ et } M \text{ no lin} \Rightarrow LX = MX = Id} = 0 = M \Delta^*$

$$\text{Cov}(\beta^*) = \sigma^2 (X^*X)^{-1} + \sigma^2 \underbrace{\Delta \Delta^*}_{\geq 0} \Rightarrow \text{Cov}(\beta^*) \geq \text{Cov}(\hat{\beta})$$

Prediction : $\hat{Y} = X \hat{\beta}$ prediction de $E(y|x)$ pour les x_i observés
 $\hat{y} = x^* \hat{\beta}$ prediction de $E(y|x)$ pour un nouvel x

Prop : $E(\hat{Y}) = X\beta$ $\text{Cov}(\hat{Y}) = \sigma^2 X (X^*X)^{-1} X^*$
 $E(\hat{y}) = x^* \beta$ $\text{Var}(\hat{y}) = \sigma^2 x^* (X^*X)^{-1} x$

Preuve : immédiat ! Pg on peut montrer que $\hat{y} = x^* \hat{\beta}$ est optimal pour la régression parmi les estimateurs linéaires $E(C^*Y) = C^*X\beta = x^*\beta$
 $\text{Var}(C^*Y) = \sigma^2 C^*C \geq \sigma^2 C^*X (X^*X)^{-1} X^*C = \sigma^2 x^* (X^*X)^{-1} x = \text{Var}(\hat{y})$

Interprétation géométrique $\hat{\beta} = \underset{\beta}{\text{argmin}} \|Y - X\tilde{\beta}\|^2$

$$\Rightarrow \hat{Y} = \underset{\tilde{Y} = X\tilde{\beta}}{\text{argmin}} \|Y - \tilde{Y}\|^2$$

$$= \underset{\tilde{Y} \in \text{Im } X}{\text{argmin}} \|Y - \tilde{Y}\|^2$$

$\Rightarrow \hat{Y}$ est la projection orthogonale de Y sur $\text{Im } X$
 (lien avec l'espace euclidien L^2 et la méthode des moindres carrés linéaires)
 $\hat{Y} = P_{\text{Im } X} Y$ où P est une matrice de projection $\perp$

Rq $\hat{Y} = X\beta = \underbrace{X(X^*X)^{-1}X^*}_{? P_{Im X}} Y$

$H = X(X^*X)^{-1}X^*$ Matrice = matrice d'ajustement = projection?

$H^* = H$

$HH = X(X^*X)^{-1}X^* X(X^*X)^{-1}X^* = X(X^*X)^{-1}X^* = H$

$H = X(X^*X)^{-1}X^* \Rightarrow Im H \subset Im X$

$\forall z \in Im X, \exists \tilde{\beta}, z = X\tilde{\beta} \Rightarrow H z = H X \tilde{\beta} = X(X^*X)^{-1}X^* X \tilde{\beta} = X \tilde{\beta} = z$

$\Rightarrow$ on retrouve par le calcul la propriété précédente.

Résidu : $\hat{\epsilon} = Y - \hat{Y}$ "estimation" des ϵ_i

Prop : $E \begin{pmatrix} \hat{Y} \\ \hat{\epsilon} \end{pmatrix} = \begin{pmatrix} X\beta \\ 0 \end{pmatrix}$ et $Cov \begin{pmatrix} \hat{Y} \\ \hat{\epsilon} \end{pmatrix} = \sigma^2 \begin{pmatrix} H = P_{Im X} & 0 \\ 0 & I - H = P_{Im X}^\perp \end{pmatrix}$

Preuve : $E(\hat{Y}) = X\beta$ a déjà été vérifié

$E(\hat{\epsilon}) = E(Y - \hat{Y}) = E(Y) - E(\hat{Y}) = 0$

$\hat{Y} = HY = P_{Im X} Y$

$\hat{\epsilon} = Y - \hat{Y} = (I - H)Y = (I - P_{Im X})Y = P_{Im X}^\perp Y$

$\begin{pmatrix} \hat{Y} \\ \hat{\epsilon} \end{pmatrix} = \begin{pmatrix} P_{Im X} \\ P_{Im X}^\perp \end{pmatrix} Y$

$\Rightarrow Cov \begin{pmatrix} \hat{Y} \\ \hat{\epsilon} \end{pmatrix} = \begin{pmatrix} P_{Im X} \\ P_{Im X}^\perp \end{pmatrix} \begin{bmatrix} \sigma^2 I_n \end{bmatrix} \begin{pmatrix} P_{Im X}^* & P_{Im X}^{*\perp} \end{pmatrix}$

$= \sigma^2 \begin{bmatrix} P_{Im X} P_{Im X}^* & P_{Im X} P_{Im X}^{*\perp} \\ P_{Im X}^\perp P_{Im X}^* & P_{Im X}^\perp P_{Im X}^{*\perp} \end{bmatrix}$

$= \sigma^2 \begin{bmatrix} P_{Im X} & 0 \\ 0 & P_{Im X}^\perp \end{bmatrix}$

Cor $E \begin{pmatrix} \hat{\beta} \\ \hat{\epsilon} \end{pmatrix} = \begin{pmatrix} \beta \\ 0 \end{pmatrix}$ et $\text{Cor} \begin{pmatrix} \hat{\beta} \\ \hat{\epsilon} \end{pmatrix} = \begin{pmatrix} (X^*X)^{-1} & 0 \\ 0 & P_{\text{Im}X^\perp} \end{pmatrix}$

Rem: Il suffit de noter que $\hat{\beta} = (X^*X)^{-1} X^* \hat{y}$.

Prop: $\hat{\sigma}^2 = \frac{\text{SCR}(\hat{\beta})}{m-p} = \frac{\|Y - X\hat{\beta}\|^2}{m-p}$ est un estimateur sans biais de σ^2

Rem: $\text{SCR}(\hat{\beta}) = \|Y - X\hat{\beta}\|^2 = \|\hat{\epsilon}\|^2$

or comme $E(\hat{\epsilon}) = 0$ $E(\|\hat{\epsilon}\|^2) = E\text{Tr}(\hat{\epsilon}\hat{\epsilon}^*)$
 $= \text{Tr}(E(\hat{\epsilon}\hat{\epsilon}^*))$

et avec $E(\hat{\epsilon}) = 0$ $E(\|\hat{\epsilon}\|^2) = \text{Tr}(\text{Cor}(\hat{\epsilon}))$

or $\text{Tr}(\text{Cor}(\hat{\epsilon})) = \text{Tr}(\sigma^2 P_{\text{Im}X^\perp})$

$= \sigma^2 (m-p)$ car $\dim \text{Im}X^\perp = m-p$

Ex: régression linéaire simple

$y_i = \beta_0 + \beta_1 x_i + \epsilon_i$

$Y = X\beta + \epsilon$

avec $X = \begin{pmatrix} 1 & x_1 \\ \vdots & \vdots \\ 1 & x_m \end{pmatrix}$ et $\beta = \begin{pmatrix} \beta_0 \\ \beta_1 \end{pmatrix}$

$X^*X = \begin{pmatrix} m & \sum x_i \\ \sum x_i & \sum x_i^2 \end{pmatrix} = m \begin{bmatrix} 1 & \bar{x} \\ \bar{x} & \overline{x^2} \end{bmatrix}$

$(X^*X)^{-1} = \frac{1}{m} \frac{1}{\overline{x^2} - (\bar{x})^2} \begin{bmatrix} \overline{x^2} & -\bar{x} \\ -\bar{x} & 1 \end{bmatrix} =$

$X^*Y = \begin{pmatrix} \sum y_i \\ \sum x_i y_i \end{pmatrix} = m \begin{pmatrix} \bar{y} \\ \overline{xy} \end{pmatrix}$

$$\hat{\beta} = (X^* X)^{-1} X^* Y = \frac{1}{\bar{x}^2 - \bar{x}^2} \begin{bmatrix} \bar{x}^2 & -\bar{x} \\ -\bar{x} & 1 \end{bmatrix} \begin{pmatrix} \bar{y} \\ \bar{x}\bar{y} \end{pmatrix}$$

$$= \frac{1}{\bar{x}^2 - \bar{x}^2} \begin{bmatrix} \bar{x}^2 \bar{y} - \bar{x} \bar{x}\bar{y} & \bar{x}^2 \bar{y} + \bar{x} (\bar{x}\bar{y}) \\ -\bar{x}\bar{y} & + \bar{x}\bar{y} \end{bmatrix}$$

$$\begin{pmatrix} \hat{\beta}_0 \\ \hat{\beta}_1 \end{pmatrix} = \begin{bmatrix} \bar{y} - \frac{\bar{x}\bar{y} - \bar{x}\bar{y}}{\bar{x}^2 - \bar{x}^2} \bar{x} \\ \frac{\bar{x}\bar{y} - \bar{x}\bar{y}}{\bar{x}^2 - \bar{x}^2} \end{bmatrix} = \begin{bmatrix} \bar{y} - \frac{\text{cov}(x, y)}{\text{var}(x)} \bar{x} \\ \frac{\text{cov}(x, y)}{\text{var}(x)} \end{bmatrix}$$

$$\rho(x, y) = \frac{\text{cov}(x, y)}{\sqrt{\text{var}(x)} \sqrt{\text{var}(y)}} = \frac{\text{cov}(x, y)}{\sigma(x) \sigma(y)}$$

$$\begin{pmatrix} \hat{\beta}_0 \\ \hat{\beta}_1 \end{pmatrix} = \begin{pmatrix} \bar{y} - \frac{\rho(x, y) \sigma(y)}{\sigma(x)} \bar{x} \\ \frac{\rho(x, y) \sigma(y)}{\sigma(x)} \end{pmatrix}$$

$$\hat{y} = \bar{y} + \frac{\rho(x, y) \sigma(y)}{\sigma(x)} (x - \bar{x}) \quad \frac{\hat{y} - \bar{y}}{\sigma(y)} = \rho(x, y) \frac{x - \bar{x}}{\sigma(x)}$$

$$\hat{y} = \bar{y} 1 - \frac{\rho(x, y) \sigma(y)}{\sigma(x)} (X - \bar{x} 1)$$

$$Y - \hat{y} = (Y - \bar{y} 1) - \frac{\rho(x, y) \sigma(y)}{\sigma(x)} (X - \bar{x} 1)$$

$$\begin{aligned} \frac{1}{n} \|Y - \hat{y}\|^2 &= \frac{1}{n} \|Y - \bar{y} 1\|^2 - 2 \frac{\rho(x, y) \sigma(y)}{\sigma(x)} \frac{1}{n} \langle Y - \bar{y} 1, X - \bar{x} 1 \rangle \\ &\quad + \frac{\rho^2(x, y) \sigma^2(y)}{\sigma^2(x)} \frac{1}{n} \|X - \bar{x} 1\|^2 \\ &= \sigma^2(y) - 2 \frac{\rho(x, y) \text{cov}(x, y) \sigma(y)}{\sigma(x)} + \rho^2(x, y) \sigma^2(y) \\ &= \sigma^2(y) (1 - \rho^2(x, y)) \end{aligned}$$

$$SCR = \|Y - \hat{Y}\|^2 = m \sigma^2(y) (1 - \rho^2(x, y))$$

$$\hat{\sigma}^2 = \frac{1}{m-2} \|Y - \hat{Y}\|^2 = \frac{m}{m-2} \sigma^2(y) (1 - \rho^2(x, y))$$

$$\text{Cov} \begin{pmatrix} \hat{\beta}_0 \\ \hat{\beta}_1 \end{pmatrix} = (X^* X)^{-1} = \frac{1}{m \text{Var}(x)} \begin{bmatrix} \overline{x^2} & -\bar{x} \\ -\bar{x} & 1 \end{bmatrix}$$

$$\Rightarrow \text{Var}(\hat{\beta}_1) = \frac{1}{m \text{Var}(x)} \leadsto x \text{ dispersé}$$

$$\text{Var}(\hat{\beta}_0) = \frac{1}{m} \left[1 + \frac{\bar{x}^2}{\overline{x^2} - \bar{x}^2} \right] \leadsto \bar{x} \text{ proche de } 0$$

Rq estimateurs non corrélés si $\bar{x} = 0$

III le modèle linéaire gaussien

on pose de $H_E = \begin{cases} E(\varepsilon) = 0 \\ \text{Cov}(\varepsilon) = \sigma^2 I_m \end{cases}$

à $\varepsilon \sim \mathcal{N}(0, \sigma^2 I_m)$

Hypothèse beaucoup plus forte $\Rightarrow$ résultats plus forts
caractérisation des lois

Modèle : $Y \sim \mathcal{N}(X\beta, \sigma^2 I_m)$

Prop : $\hat{\beta} = (X^* X)^{-1} X^* Y$, $\hat{Y} = X \hat{\beta} = P_{I_m X} Y$

et $\hat{\varepsilon} = Y - \hat{Y} = P_{I_m X^\perp} Y$

• $\hat{\beta} \sim \mathcal{N}(\beta, \sigma^2 (X^* X)^{-1})$, $\hat{Y} \sim \mathcal{N}(X\beta, \sigma^2 P_{I_m X})$

et $\hat{\varepsilon} \sim \mathcal{N}(0, \sigma^2 P_{I_m X^\perp})$

• $(\hat{\beta}, \hat{Y}) \perp \hat{\varepsilon}$

Preuve : le premier point est un rappel

le second utilise le fait que l'image par une application linéaire d'un vecteur gaussien est un vecteur gaussien et que celui-ci est entièrement caractérisé par sa moyenne et sa covariance que l'on a déjà calculé

le troisième point s'obtient en notant que $\begin{pmatrix} \hat{\beta} \\ \hat{\varepsilon} \end{pmatrix}$ est un vecteur gaussien dont la covariance a une structure par bloc $\begin{pmatrix} \sigma^2 I_{m-p} & 0 \\ 0 & \sigma^2 I_p \end{pmatrix}$

Prop $\|\hat{\varepsilon}\|^2 \sim \sigma^2 \chi^2(m-p)$

Preuve: $\hat{\varepsilon} \sim \mathcal{N}(0, P_{\perp} \sigma^2)$

donc $\exists B$ matrice orthogonale tel que $B P_{\perp} B^* = \begin{pmatrix} I_{m-p} & 0 \\ 0 & 0 \end{pmatrix}$

$\Rightarrow B \hat{\varepsilon} \sim \mathcal{N}\left(0, \sigma^2 \begin{pmatrix} I_{m-p} & 0 \\ 0 & 0 \end{pmatrix}\right)$

$\Rightarrow (B \hat{\varepsilon})_i = \sigma \varepsilon_i, 1 \leq i \leq m-p$ avec ε_i iid

$\|\beta \hat{\varepsilon}\|^2 = \sum_{i=1}^{m-p} (B \hat{\varepsilon})_i^2 = \sum_{i=1}^{m-p} \sigma^2 \varepsilon_i^2$

$\sim \sigma^2 \chi^2(m-p)$

Prop: propriété générale des gaussiennes

$\varepsilon_1 \sim \mathcal{N}(0, P)$ et $\varepsilon_2 \sim \mathcal{N}(0, Q)$ avec P et Q deux matrices orthogonales orthogonales.

$\Rightarrow \|\varepsilon_1\|^2 \sim \chi^2(\text{rang } P) \quad \|\varepsilon_2\|^2 \sim \chi^2(\text{rang } Q)$

et $\|\varepsilon_1\|^2 \perp \|\varepsilon_2\|^2$

■ Lien avec le maximum de vraisemblance

$\tilde{Y} \sim \mathcal{N}(X\beta, \sigma^2 I_m)$

$\Rightarrow dP(\tilde{Y}=Y; \beta, \sigma^2) = \frac{1}{(2\pi\sigma^2)^{m/2}} e^{-\frac{\|Y-X\beta\|^2}{2\sigma^2}}$

$\Rightarrow$ La log vraisemblance des observations est donnée

$L_m(\beta, \sigma^2) = -\frac{m}{2} [\log 2\pi + \log \sigma^2] - \frac{1}{2\sigma^2} \|Y-X\beta\|^2$

$\Rightarrow$ Estimateur du maximum de vraisemblance

$\frac{\partial L_m}{\partial \beta} = X^* (Y - X\beta) \Rightarrow \hat{\beta}^{MV} = (X^* X)^{-1} X^* Y = \hat{\beta}$

$\frac{\partial L_m}{\partial \sigma^2} = -\frac{m}{2\sigma^2} + \frac{1}{2\sigma^4} \|Y-X\beta\|^2$

$\Rightarrow \hat{\sigma}^{2MV} = \frac{\|Y-X\hat{\beta}\|^2}{m} = \frac{m-p}{m} \hat{\sigma}^2$

liée mais croissant en espérance!

Prop Parmi les estimateurs sans biais, $\hat{\beta}$ est l'estimateur de moindre variance (BLUE Best Unbiased Estimator)

Preuve La matrice de Fisher est donnée (si β est régulier)

$$\begin{aligned} J(\beta) &= E \left(- \left(\frac{\partial \ln(\beta \sigma^2)}{\partial \beta} \right)^* \left(\frac{\partial \ln(\beta \sigma^2)}{\partial \beta} \right) \right) \\ &= E \left(- \frac{\partial^2 \ln(\beta \sigma^2)}{\partial \beta^2} \right) \stackrel{\text{calcul déjà fait pour minimiser}}{=} E \left[\frac{1}{\sigma^2} X^* X \right] / SSR(\beta) \\ &= \frac{1}{\sigma^2} (X^* X) \end{aligned}$$

Le théorème de Gauss-Markov indique qu'un estimateur sans biais T de β doit satisfaire $\text{Var}(T) \geq (J(\beta))^{-1} = \sigma^2 (X^* X)^{-1}$.
L'estimateur des moindres carrés atteint cette borne, il est donc optimal.

Prédiction $\hat{y} = X^* \hat{\beta} \sim N(X^* \beta, X^* (X^* X)^{-1} X)$

IV Tests du modèle linéaire gaussien $\rightarrow$ Voir plus loin

V Analyse asymptotique du modèle linéaire

$$y_i = x_i^* \beta + \varepsilon_i \quad 1 \leq i \leq n$$

Sur H_0 que se passe-t-il quand n grandit? $X_n = \begin{pmatrix} x_1^* \\ \vdots \\ x_n^* \end{pmatrix}$ joue un rôle important

Prop Sous H_0 , $\hat{\beta}_n \xrightarrow{P} \beta$ dès que $\text{tr}(X_n^* X_n)^{-1} \rightarrow 0$

Preuve $E(\hat{\beta}_n) = \beta$ et $\text{Var}(\hat{\beta}_n) = (X_n^* X_n)^{-1}$

$$E \|\hat{\beta}_n - \beta\|_{\text{tr}}^2 = \text{tr}(X_n^* X_n)^{-1} \rightarrow 0$$

Cor Sous H_0 , $\text{tr}(X_n^* X_n)^{-1} \rightarrow 0 \Rightarrow \hat{\beta}_n \xrightarrow{P} \beta$

Pr $\text{tr}(X_n^* X_n)^{-1} \rightarrow 0 \Rightarrow (X_n^* X_n)^{-1} \rightarrow 0$ la matrice de covariance

Ex Si X_i iid de loi Z telle que $Q = E(Z^* Z)$ d.p

$$\Rightarrow \frac{1}{n} X_n^* X_n \rightarrow Q \quad \text{par la loi faible des grands nombres}$$

$$\Rightarrow n (X_n^* X_n)^{-1} \rightarrow Q^{-1}$$

$$\Rightarrow (X_n^* X_n)^{-1} \rightarrow 0$$

• Sous $H_0 + \varepsilon_i$ indépendant + condition (celiques)

• $\frac{1}{n} X_n^* X_n \rightarrow Q$ définie positive ("Loynque")

H_{ex}^+ • $\sum \|x_i\|^{2+\delta} |\varepsilon_i|^{2+\delta} = o(n^{1+\delta/2})$ ("Technique")

(OK si ε_i iid $E(|\varepsilon_i|^{2+\delta}) < +\infty$ et $\sum \|x_i\|^{2+\delta} = o(n^{1+\delta/2})$)

(OK si $\|x_i\| < C$ ou $E(\|x_i\|^{2+\delta}) < +\infty$)

Prop Sous H_{ex}^+ , $\sqrt{n}(\hat{\beta}_n - \beta) \xrightarrow{L} \mathcal{N}(0, \sigma^2 Q^{-1})$

Preuve

• $E(\hat{\beta}_n) = \beta$ et $\text{Cov}(\hat{\beta}_n) = \sigma^2 (X_n^* X_n)^{-1}$

dnc $E(\sqrt{n}(\hat{\beta}_n - \beta)) = 0$ et $\text{Cov}(\sqrt{n}(\hat{\beta}_n - \beta)) = \sigma^2 \left(\frac{1}{n} X_n^* X_n \right)^{-1}$

• $\sqrt{n}(\hat{\beta}_n - \beta) = \sqrt{n} \left[(X_n^* X_n)^{-1} X_n^* Y - \beta \right]$

$= \sqrt{n} (X_n^* X_n)^{-1} [X_n^* Y - X_n^* X_n \beta]$

$= \sqrt{n} (X_n^* X_n)^{-1} X_n^* [Y - X_n \beta]$

$= \sqrt{n} (X_n^* X_n)^{-1} X_n^* \varepsilon_n$

$= \left(\frac{1}{n} X_n^* X_n \right)^{-1} \underbrace{\frac{1}{\sqrt{n}} X_n^* \varepsilon_n}_{Z_n}$

$Q^* Z_n = \frac{1}{\sqrt{n}} Q^* X_n^* \varepsilon_n = \frac{1}{\sqrt{n}} (X_n a)^* \varepsilon_n$

$= \frac{1}{\sqrt{n}} \sum_{i=1}^n \underbrace{(x_i^* a)}_{z_i} \varepsilon_i \times \frac{1}{\sqrt{n}}$

• $\beta_i \perp$

• $E(\varepsilon_i) = 0$

$S_n^2 = \sum_{i=1}^n E(|\varepsilon_i|^2)$
 $K_n = \sum_{i=1}^n E(\| \beta_i \|^{2+\delta})$
 $S_n^2 = \sum_{i=1}^n \frac{1}{n} (x_i^* a)^2 E(\varepsilon_i)^2 = \sigma^2 \frac{1}{n} a^* X_n^* X_n a \rightarrow \sigma^2 a^* Q a > 0$

$K_n = \sum \frac{1}{n^{1+\delta/2}} |k_i^* a|^{2+\delta} E|\varepsilon_i|^{2+\delta} = o(1) \Rightarrow \frac{K_n}{(S_n^2)^{2+\delta}} \rightarrow 0$

TA de Lagrange s'applique

$$\frac{1}{S_n} \sum_{i=1}^n \beta_i \xrightarrow{\mathcal{L}} \mathcal{N}(0,1)$$

$$\frac{1}{S_n} \frac{a^*}{\sqrt{n}} Z_n \xrightarrow{\mathcal{L}} \mathcal{N}(0,1)$$

$$S_n^2 \rightarrow \sigma^2 \quad a^* Q_n a \rightarrow a^* a^* Q a$$

$$\Rightarrow a^* Z_n \xrightarrow{\mathcal{L}} \mathcal{N}(0, \sigma^2 a^* Q a)$$

$$\Rightarrow Z_n \xrightarrow{\mathcal{L}} \mathcal{N}(0, \sigma^2 Q)$$

$$\text{or } \sqrt{n} (\hat{\beta}_n - \beta) = Q_n^{-1} Z_n$$

$$\text{or } Q_n^{-1} \rightarrow Q^{-1}$$

$$\Rightarrow \sqrt{n} (\hat{\beta}_n - \beta) \rightarrow \mathcal{N}(0, \sigma^2 Q^{-1})$$

Rq En pratique, on utilise Q_n au lieu de Q

$$\sqrt{n} Q_n^{1/2} (\hat{\beta}_n - \beta) \xrightarrow{\mathcal{L}} \mathcal{N}(0, \sigma^2 I_p)$$

■ Convergence pour $\hat{\sigma}^2$

Prop : Sous $H_{\varepsilon, X}^1$, $\hat{\sigma}^2 \xrightarrow{P} \sigma^2$
+ ε_i i.i.d

$$\begin{aligned} \text{Preuve : } \hat{\sigma}^2 &= \frac{\|Y - \hat{Y}\|^2}{n-p} = \frac{n}{n-p} \frac{\|(I - P_{\text{Inx}})Y\|^2}{n} \\ &= \frac{n}{n-p} \frac{\|(I - P_{\text{Inx}})\varepsilon\|^2}{n} \\ &= \frac{n}{n-p} \frac{\varepsilon^* (I - P_{\text{Inx}}) \varepsilon}{n} \\ &= \frac{n}{n-p} \left[\frac{\varepsilon^* \varepsilon}{n} - \frac{\varepsilon^* P_{\text{Inx}} \varepsilon}{n} \right] \\ &\quad \downarrow \quad \quad \downarrow P \\ &\quad 1 \quad \quad \hat{\sigma}^2 \quad \quad \text{(loi des grands nombres)} \end{aligned}$$

$$\begin{aligned}
\frac{E^* P_{Im \times n} E}{n} &= \frac{E^* X_m (X_m^* X_m)^{-1} X_m^* E}{n} \\
&= \frac{1}{\sqrt{n}} (X_m^* E)^* (X_m^* X_m)^{-1} \frac{1}{\sqrt{n}} (X_m^* E) \\
&= Z_m (X_m^* X_m)^{-1} Z_m \\
&= \frac{1}{n} Z_m Q_m^{-1} Z_m \\
&= \frac{1}{n} \underbrace{\| Q_m^{-1/2} Z_m \|_2^2}_{\substack{\sim \mathcal{N}(0, \sigma^2 I_p)}} \xrightarrow{P_1} 0
\end{aligned}$$

Cor Si $E(|\epsilon_i|^4) < +\infty$,
 $\sqrt{n}(\hat{\sigma}^2 - \sigma^2) \rightarrow \mathcal{N}(0, \text{Var}(|\epsilon_i|^2))$

Preuve

$$\begin{aligned}
\sqrt{n}(\hat{\sigma}^2 - \sigma^2) &= \sqrt{n} \left[\frac{E^* E}{n} - \sigma^2 \right] + \frac{p}{\sqrt{n}} \frac{E^* E}{n} + \sqrt{n} \frac{E^* P_{Im \times n} E}{n-p} \\
&\quad \downarrow \mathcal{L} \quad \downarrow \frac{p \sigma^2}{\sqrt{n}} \downarrow 0 \quad \downarrow \sqrt{n} \frac{p}{n-p} \xrightarrow{P_2} 0 \\
&\quad \mathcal{N}(0, \text{Var}(|\epsilon_i|^2))
\end{aligned}$$

IV Test dans le modèle linéaire gaussien

○ Tests simples

■ Test de Student

(H₀) $\beta_q = a$ contre (H₁) $\beta_q \neq 0$

Prop : sous H₀ $\frac{\hat{\beta}_q - a}{\hat{\sigma} \sqrt{(X^* X)^{-1}_{qq}}} \sim T_{(n-p)}$ Student

Preuve sous H₀ $\beta_q \sim \mathcal{N}(a, \sigma^2 (X^* X)^{-1}_{qq})$ et $\hat{\sigma}^2 \sim \sigma^2 \chi^2_{(n-p)}$

Rq P-value

$$\begin{aligned}
&2 Q_{T(n-p)}^{-1} \left(\left| \frac{\hat{\beta}_q - a}{\hat{\sigma} \sqrt{(X^* X)^{-1}_{qq}}} \right| \right) \\
&= 2(1 - F_{T(n-p)} \left(\left| \frac{\hat{\beta}_q - a}{\hat{\sigma} \sqrt{(X^* X)^{-1}_{qq}}} \right| \right)) \\
&\approx \alpha \text{ fonction de } \left(\left| \frac{\hat{\beta}_q - a}{\hat{\sigma} \sqrt{(X^* X)^{-1}_{qq}}} \right| \right)
\end{aligned}$$

la zone de rejet

Cor : Test $\left| \frac{\hat{\beta}_q - a}{\hat{\sigma} \sqrt{(X^* X)^{-1}_{qq}}} \right| > Q_{T(n-p)}^{-1} \left(1 - \frac{\alpha}{2} \right) = \tilde{Q}_{T(n-p)} \left(\frac{\alpha}{2} \right)$

Rq $\Leftrightarrow$ est de niveau α ou IC $F_T(Q) = 1 - \alpha$ $\tilde{Q}(\alpha) = Q(1 - \alpha)$

o Test sur une contrainte linéaire (Généralisation)

$$H_0: b^* \beta = a \quad \text{contre} \quad H_1: b^* \beta \neq a$$

Prq sous $H_0: \frac{b^* \hat{\beta} - a}{\hat{\sigma} \sqrt{b^* (X^* X)^{-1} b}} \sim T(m-p)$

$\Rightarrow$ Test similaire au précédent

o Test sur σ^2

$$H_0: \sigma = \sigma_0 \quad \text{contre} \quad H_1: \sigma \neq \sigma_0$$

Prq sous $H_0: (m-p) \frac{\hat{\sigma}^2}{\sigma_0^2} = \frac{1}{\sigma_0^2} \sum |y_i - x_i^* \hat{\beta}|^2 \sim \chi^2(m-p)$

$$= \frac{SSR(\hat{\beta})}{\sigma_0^2}$$

$\Rightarrow$

$$\underline{\text{Test}} \leq \frac{SSR(\hat{\beta})}{\sigma_0^2} \geq Q_{\chi^2(m-p)} \left(1 - \frac{\alpha}{2}\right)$$

$$\text{ou} \quad \frac{SSR(\hat{\beta})}{\sigma^2} \leq Q_{\chi^2(m-p)} \left(\frac{\alpha}{2}\right)$$

$$p \text{ values} = \min \left(2 \left(1 - F_{\chi^2} \left(\frac{SSR(\hat{\beta})}{\sigma^2} \right) \right), 2 F_{\chi^2} \left(\frac{SSR(\hat{\beta})}{\sigma^2} \right) \right)$$

$\Rightarrow$ Intervalle de confiance

o Test de Fisher

$$\mathcal{M}_1: Y \sim \mathcal{N}(X\beta, \sigma^2 I_m) \quad \beta \in \mathbb{R}^p$$

$$\mathcal{M}_0: Y \sim \mathcal{N}(Z\gamma, \sigma^2 I_m) \quad \gamma \in \mathbb{R}^q \quad q < p$$

on dit que $\mathcal{M}_0$ est un sous modèle de $\mathcal{M}_1$

(car que les modèles sont emboîtés)

$$\text{si} \quad \text{Im } Z \subset \text{Im } X$$

Ex : $\beta = H\delta$ cf analyse de la variance

- imposition de $r = p - q$ contraintes linéaires
 $\leadsto$ sous-espace de dimension q à paramètres.

Prédiction par ordre de complexité

$$\begin{array}{ccc} Y & \hat{Y}^{\alpha_1} & \hat{Y}^{\alpha_0} \\ \parallel & & \parallel \\ P_{\text{Im } X} Y & & P_{\text{Im } Z} Y \end{array}$$

Comme $\text{Im } Z \subset \text{Im } X$ on en déduit

$$\text{que } I - P_{\text{Im } X} \perp P_{\text{Im } X} - P_{\text{Im } Z}$$

Généralisation gaussienne $\begin{array}{ccc} (I - P_{\text{Im } X}) Y & \parallel & (P_{\text{Im } X} - P_{\text{Im } Z}) Y \\ Y - \hat{Y}^{\alpha_1} & \parallel & \hat{Y}^{\alpha_1} - \hat{Y}^{\alpha_0} \end{array}$

Sous H_0 : α_0 est valide

$$Y - \hat{Y}^{\alpha_1} \sim \mathcal{N}(0, \sigma^2 P_{\text{Im } X}^\perp) \quad (\text{c'est un échantillon de } \mathcal{N})$$

$$\hat{Y}^{\alpha_1} - \hat{Y}^{\alpha_0} \sim \mathcal{N}(0, \sigma^2 P_{\text{Com } \text{Im } Z \text{ des } \text{Im } X})$$

$$\Rightarrow \begin{array}{l} \|Y - \hat{Y}^{\alpha_1}\|^2 \sim \sigma^2 X^2(m-p) \\ \|\hat{Y}^{\alpha_1} - \hat{Y}^{\alpha_0}\|^2 \sim \sigma^2 X^2(p-q) \end{array} + \parallel$$

$$\Rightarrow \frac{\|\hat{Y}^{\alpha_1} - \hat{Y}^{\alpha_0}\|^2 / (p-q)}{\|Y - \hat{Y}^{\alpha_1}\|^2 / (m-q)} \sim T(p-q, m-p)$$

Puis
$$\frac{[SSR(\alpha_0) - SSR(\alpha_1)] / (p-q)}{SSR(\alpha_1) / (m-q)} \sim T(p-q, m-q)$$

Here $\|Y - \hat{Y}^{\alpha_1}\|^2 = SSR(\alpha_1)$

$$\|Y - \hat{Y}^{\alpha_0}\|^2 = SSR(\alpha_0) = \|\hat{Y} - Y^{\alpha_1}\|^2 + \|Y^{\alpha_1} - Y^{\alpha_0}\|^2$$

$$\Rightarrow \|\hat{Y}^{\alpha_1} - \hat{Y}^{\alpha_0}\|^2 = SSR(\alpha_0) - SSR(\alpha_1)$$

⇒ Test de Fischer.

0 Test de Wald

$$\mathcal{N}(X\beta, \sigma^2 I_m) \text{ avec } \beta \in \mathbb{R}^p$$

soit C m $p \times q$ de rang $(p-q)$ r.f. $p-q$ colonnes avec $p < q < p$
et $c \in \mathbb{R}^{p-q}$

$$H_0: C\beta = c \quad \text{contre} \quad H_1: C\beta \neq c$$

$$\text{Sous } H_0 \quad C\hat{\beta} - c \sim \mathcal{N}(0, \sigma^2 C (X^*X)^{-1} C^*)$$

$$\underline{\text{Lem}} \quad Z \sim \mathcal{N}_m(0, \Gamma) \text{ avec } \Gamma \text{ inversible} \\ \Rightarrow Z^* \Gamma^{-1} Z \sim \chi_m^2$$

$$\underline{\text{Preuve}} \quad \Gamma^{-1} = P D P \quad \text{avec } \Gamma \text{ inversible e matrix de cov} \\ = P D^{1/2} D^{1/2} P \quad D \geq 0$$

$$Z^* \Gamma^{-1} Z = \| D^{1/2} P Z \|^2$$

$$\text{or } D^{1/2} P Z \sim \mathcal{N}(0, \underbrace{D^{1/2} P \Gamma^{-1} P D^{1/2}}_{D^{1/2} P P^{-1} D^{-1} P^{-1} P D^{1/2}})$$

$$\underline{\text{Prop}} \quad \frac{[(C\hat{\beta} - c)^* [C (X^*X)^{-1} C^*]^{-1} C(\hat{\beta} - c)] / (p-q)}{\hat{\sigma}^2} \sim T(p-q, m-p)$$

$$\underline{\text{Preuve}} \quad \text{Lemme + iid de } \hat{\beta} \text{ et } \hat{\sigma}^2$$

$$\underline{\text{Cor}} \quad C = Id \quad \hat{\beta} = \beta \\ \frac{[(\hat{\beta} - \beta)^* (X^*X) (\hat{\beta} - \beta)] / p}{\hat{\sigma}^2} \sim T(p, m-p)$$

⇒ Ellipsoïde de confiance

Ex 11 Test de rupture : analyse de la course vers l'union

Période 1 R^2 $Y_1 \sim \mathcal{N}(X_1 \beta_1, \sigma_1^2 I_n)$ $\beta_1 \in \mathbb{R}^p$ m_1 obs

Période 2 $Y_2 \sim \mathcal{N}(X_2 \beta_2, \sigma_2^2 I_m)$

Test d'essai de rupture

• Test $\sigma_1^2 = \sigma_2^2$

$$\frac{SCR(\hat{\beta}_1) / (m_1 - p)}{SCR(\hat{\beta}_2) / (m_2 - p)} \sim F(m_1 - p, m_2 - p)$$

• Modèle "complet"

Test $\beta_1 = \beta_2$ par la méthode de Fisher

• Test de Fisher du modèle

$H_0: E(Y) = \beta_0 \beta$

$H_1: E(Y) = X \beta$

Sous H_0 :
$$\frac{[SCR(H_0) - SCR(H_1)] / p - 1}{SCR(H_1) / (m - p)} \sim F(p - 1, m - p)$$

$$= \frac{\|X\hat{\beta} - \bar{Y}\|^2 / p - 1}{\|X\hat{\beta} - Y\|^2 / (m - p)}$$

• Coefficient $R^2 = \frac{\|\hat{Y} - \bar{Y}\|^2}{\|Y - \bar{Y}\|^2} \leq 1$

mesure d'ajustement de $\hat{Y}$

$$R^2 = 1 - \frac{\|Y - \hat{Y}\|^2}{\|Y - \bar{Y}\|^2}$$

on ma utilisé $\langle \hat{Y} - \bar{Y}, Y - \hat{Y} \rangle = 0$

on a également $\langle \hat{Y} - \bar{Y}, Y - \bar{Y} \rangle = \|\hat{Y} - \bar{Y}\|^2$

$$\rightarrow R^2 = \frac{\langle \hat{Y} - \bar{Y}, Y - \bar{Y} \rangle^2}{\|\hat{Y} - \bar{Y}\|^2 \|Y - \bar{Y}\|^2} = \rho(Y, \hat{Y}_e)^2$$

on vérifie que Fisher $\Omega_0 \subset \Omega_1$

$$F = \frac{n - \dim(\Omega_1)}{\dim(\Omega_1) - \dim(\Omega_0)} \frac{R_{\Omega_0}^2 - R_{\Omega_1}^2}{1 - R_{\Omega_1}^2}$$

VI Moindres carrés généralisés

$$H'_\varepsilon \begin{cases} E(\varepsilon) = 0 \\ \text{Cov}(\varepsilon) = \sigma^2 \Sigma \quad \text{avec } \Sigma \neq c I_n \end{cases}$$

2 cas classiques : $\text{Cov}(\varepsilon) = \text{Diag}(\sigma_i^2)$

Hétéroscastichité

• $\text{Cov}(\varepsilon_i, \varepsilon_j) \neq 0 \sim \varepsilon$ processus l'ombré

■ Cas Σ connu

on pose $\Sigma^{1/2}$ une racine de Σ qui existe

$$\Sigma = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix} \Rightarrow \begin{aligned} \Sigma^{1/2} &= \begin{pmatrix} 0 & 0 & 1/2 \\ 0 & 0 & 1/2 \\ 0 & 0 & 0 \end{pmatrix} \text{ orienté} \\ \Sigma^{-1/2} &= \begin{pmatrix} 0 & 0 & -1/2 \\ 0 & 0 & -1/2 \\ 0 & 0 & 0 \end{pmatrix} \end{aligned}$$

$$\tilde{y} = \Sigma^{-1/2} y$$

$$\Rightarrow = \Sigma^{-1/2} X \beta + \underbrace{\Sigma^{-1/2} \varepsilon}_{\hat{\varepsilon}}$$

$$\Rightarrow E(\hat{\varepsilon}) = 0$$

$$\text{Cov}(\hat{\varepsilon}) = \sigma^2 \text{Id}$$

$\Rightarrow$ on réécrit comme au cas ordinaire

$$\Rightarrow \hat{\beta} = (X^* \tilde{\Sigma}^{-1} X)^{-1} X^* \tilde{\Sigma}^{-1} \tilde{y}$$

Interprétation géométrique : $\Sigma^{-1} \Rightarrow$ forme quadratique définie positive
 $\Rightarrow$ norme sur $\mathbb{R}^n$
 $\|Y\|_{\Sigma^{-1}}^2 = Y^T \Sigma^{-1} Y$

$$\begin{aligned} \beta &= \text{argmin} \|\hat{Y} - \hat{X}\beta\|^2 \\ &= \text{argmin} \|\Sigma^{-1/2}(Y - X\beta)\|^2 \\ &= \text{argmin} \|Y - X\beta\|_{\Sigma^{-1}}^2 \end{aligned}$$

$\Rightarrow$ Projection Σ^{-1} orthogonale !

Prop : $E(\hat{\beta}) = \beta \quad \text{Var}(\hat{\beta}) = (X^T \Sigma^{-1} X)^{-1}$

• (Gauss-Markov) $\hat{\beta}$ est le meilleur estimateur linéaire sans biais

• $\hat{Y} = X\hat{\beta}$ est la projet Σ^{-1} orthogonal de Y sur $\text{Im } X$

• $\hat{\sigma}^2 = \frac{\|\hat{Y} - X\hat{\beta}\|^2}{n-p} = \frac{\|Y - X\hat{\beta}\|_{\Sigma^{-1}}^2}{n-p}$ estime sans biais σ^2
 $\left(\neq \frac{\|Y - X\hat{\beta}\|_Z^2}{n-p} \right)$

Preuve : immédiat

Prop : • Si Y est gaussien : $\hat{\beta}$ est gaussien $\|\hat{\sigma}^2\| \sim \sigma^2 \frac{\chi^2_{(n-p)}}{n-p}$

• Pour β coincide avec le maximum de vraisemblance

• $\hat{\beta}$ est optimal parmi les estimateurs sans biais.

Preuve $L_n(Y; \beta, \sigma^2) = (2\pi\sigma^2)^{-n/2} \det(\Sigma)^{-1/2} \exp\left(-(Y - X\beta)^T \Sigma^{-1} (Y - X\beta)\right)$

Arg : Pour calculer $\hat{\beta}$, on peut soit par $\hat{Y}$ et le calcul de $\Sigma^{-1/2}$
 • soit par la forme et le calcul de Σ
 $\Rightarrow$ calcul les SCR par des formules $\neq$

■ Cas Σ inconnu

$\leadsto$ on utilise les moindres carrés ordinaires:

$$\hat{\beta} = \arg\min \|\mathbf{Y} - \mathbf{X}\beta\|_2^2$$

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

Prop $E(\hat{\beta}) = \beta$

$$\text{Var}(\hat{\beta}) = \sigma^2 (\mathbf{X}^* \mathbf{X})^{-1} \mathbf{X}^* \Sigma \mathbf{X} (\mathbf{X}^* \mathbf{X})^{-1}$$

⚠ on perd tout le reste ⚠

Prop si $\lambda_M(\Sigma_n) = o(1)$

$$\lambda_m(\mathbf{X}^* \mathbf{X}) \xrightarrow{n \rightarrow \infty} +\infty \quad \left[\lambda_m[(\mathbf{X}^* \mathbf{X})^{-1}] = o(1) \right]$$

$$\Rightarrow \text{Tr}((\mathbf{X}^* \mathbf{X})^{-1} \mathbf{X}^* \Sigma \mathbf{X} (\mathbf{X}^* \mathbf{X})^{-1}) \rightarrow 0$$

$\Rightarrow \hat{\beta}$ est asymptotiquement convergent

Preuve: $\text{Tr}((\mathbf{X}^* \mathbf{X})^{-1} \mathbf{X}^* \Sigma \mathbf{X} (\mathbf{X}^* \mathbf{X})^{-1})$

$$= \text{Tr}(\Sigma \mathbf{X} (\mathbf{X}^* \mathbf{X})^{-2} \mathbf{X}^*)$$

$$\begin{aligned} * & \leq \lambda_M(\Sigma) \text{Tr}(\mathbf{X} (\mathbf{X}^* \mathbf{X})^{-2} \mathbf{X}^*) \leq \lambda_M(\Sigma) \text{tr}((\mathbf{X}^* \mathbf{X})^{-1}) \\ & \leq \lambda_M(\Sigma) p \lambda_m((\mathbf{X}^* \mathbf{X})^{-1}) \rightarrow 0 \end{aligned}$$

* $\Sigma = \mathbf{O}^* \mathbf{D} \mathbf{O}$ avec $\mathbf{O}$ orthogonale

$$\text{Tr}(\Sigma \mathbf{V}) = \text{Tr}(\mathbf{O}^* \mathbf{D} \mathbf{O} \mathbf{V})$$

$$= \text{Tr}(\mathbf{D} \mathbf{O} \mathbf{V} \mathbf{O}^*)$$

$$\leq (\sum_i D_{ii}) \text{tr}(\mathbf{O} \mathbf{V} \mathbf{O}^*)$$

$$\leq \lambda_M(\Sigma) \text{tr}(\mathbf{V})$$

- Estimation de Σ impossible : n^2 paramètres
mais estimation "possible" si Σ a une forme paramétrique simple...

VII Etude des résidus

1/ Tests

■ résidu $\hat{\varepsilon} = Y - \hat{Y} \approx \varepsilon = Y - E(Y)$

$$E(\hat{\varepsilon}) = 0$$

$$\text{var}(\hat{\varepsilon}) = \sigma^2(1-H)$$

$$E(\varepsilon) = 0$$

$$\text{var}(\varepsilon) = \sigma^2 I$$

■ $\hookrightarrow \text{Var } \hat{\varepsilon}_i = \sigma^2(1-h_{ii})$

$\hookrightarrow \text{Var}(\varepsilon_i) = \sigma^2$

- Modification des résidus pour les rendre de même vari

$$\tilde{\varepsilon}_i = \frac{\hat{\varepsilon}_i}{\sqrt{1-h_{ii}}}$$

$$\tilde{\varepsilon}_i = \frac{\hat{\varepsilon}_i}{\sigma \sqrt{1-h_{ii}}}$$

$\Rightarrow$ Standardisation en utilisant l'estimation de la variance

$$t_i = \frac{\hat{\varepsilon}_i}{\hat{\sigma} \sqrt{1-h_{ii}}}$$

ou $t_i^* = \frac{\hat{\varepsilon}_i}{\hat{\sigma}_{(i)} \sqrt{1-h_{ii}}}$

où $\hat{\sigma}_{(i)}$ estimation du σ sur le modèle où l'on enlève la i -ème observation

R $t_i \approx \mathcal{N}(0,1)$

$t_i^* \sim F(n-1-p)$

P pour calculer $\hat{\sigma}_{(i)}$ $\Rightarrow$ Formule "magique"

$$\frac{\hat{\sigma}^2}{\hat{\sigma}_{(i)}^2} = \frac{n-p}{n-1-p-t_i^2}$$

$$\Rightarrow t_i^* = t_i \sqrt{\frac{n-p}{n-1-p-t_i^2}}$$

⚠ Pas de décorrélation entre les $\hat{\epsilon}_i, \hat{\epsilon}_i^2, \hat{\epsilon}_i^3$. ⚠

En pratique, on fait les tests sur $\hat{\epsilon}_i$ ou $\hat{\epsilon}_i^2$.

⇒ Test de symétrie

- Wilcoxon

basé sur $\sum_i \text{sign}(\hat{\epsilon}_{ci})$

⇒ Test

⇒ Test de symétrie / indépendance

- Wald - Wolfowitz Test des séquences

$\text{sign}(\hat{\epsilon}_{ci}) = ++--++-+-$

⇒ on compte le nombre de "suite de m signes"

⇒ Test d'indépendance

- Durbin - Watson

→ Estimation de ρ dans un modèle AR(1)

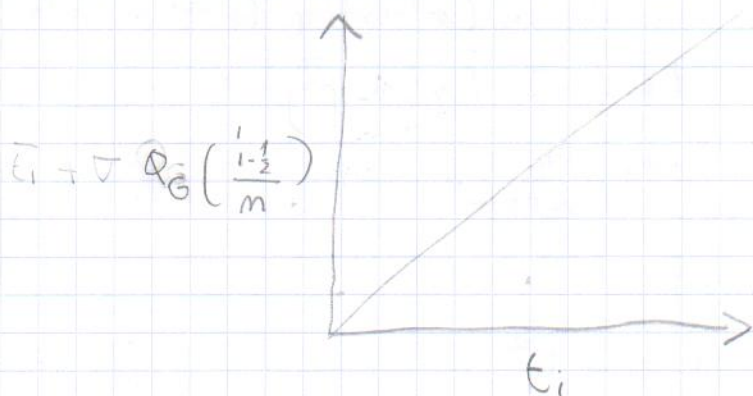
$$\hat{\epsilon}_i = \rho \hat{\epsilon}_{i-1} + \sqrt{1-\rho^2} \omega_i$$

+ Test de $\rho = 0$

⇒ Test de normalité

- Pearson : Test du χ^2 sur des boîtes équiprobables

- QQ-plot de Shapiro - Wilks



Shapiro - Wilks

, test sur la corrélation

entre $(\hat{\epsilon}_i)$ et $(Q_0(\frac{i-1/2}{n}))$

• Lilliefors (Kolmogorov-Smirnov)

$$D_+ = \max \left(\frac{i}{n} - F_G \left(\frac{x_{(i)} - \bar{x}}{\hat{\sigma}} \right) \right)$$

$$D_- = \max \left(F_G \left(\frac{x_{(i)} - \bar{x}}{\hat{\sigma}} \right) - \frac{i-1}{n} \right)$$

$$D = \max(D_+, D_-)$$

$$\Rightarrow \|F_n - F_0\|_\infty : \text{Kolmogorov-Smirnov}$$

• Cramer-von-Mises

$$\sum_{i=1}^n \left(F_G \left(\frac{x_{(i)} - \bar{x}}{\hat{\sigma}} \right) - \frac{i - \frac{1}{2}}{n} \right)^2$$

• Andersen-Darling

$$\sum_{i=1}^{n/2} (2i-1) \left[\ln \left(F_G \left(\frac{x_{(i)} - \bar{x}}{\hat{\sigma}} \right) \right) + \ln \left[1 - F_G \left(\frac{x_{(n-i+1)} - \bar{x}}{\hat{\sigma}} \right) \right] \right]$$

plus de poids que le test de Cramer-von-Mises

• Shapiro-Wilk

$$\frac{1}{\hat{\sigma}^2} \sum_{i=1}^{n/2} a_i \left(t_{[n-i+1]} - t_{(i)} \right)^2 \quad (\text{la règle de Shapiro-Wilk})$$

• Test de d'Agostino / Test de Jarque-Bera

Tests liés aux relations entre les moments des gaussiennes.

$\Rightarrow$ Test de sphéricité

2/ Diagnostics

Tout ce qui comportement "bizarre" (différent) par les données

■ leverage = "lever"

$$h_{ii} = x_i^* (X^* X)^{-1} x_i = X^* (X X)^{-1} X = H$$

≈ distance au centre du nuage des (x_i)
mesurée par la distance de Mahalanobis associée au nuage de points.

$$\hat{y}_i = h_{ii} y_i + \sum_{j \neq i} h_{ij} y_j$$

Prop : $\sum h_{ii} = p$

Preuve $\sum h_{ii} = \text{Tr}(H) = p$

Heuristique : si tous les points ont aussi importance, $h_{ii} \approx \frac{p}{n}$

Points à étudier si $h_i > C \frac{p}{n}$ $C = 2$ ou 3

■ Résidu standardisé $t_i = \frac{\hat{\epsilon}_i}{\hat{\sigma} \sqrt{1-h_{ii}}}$

Heuristique : si t_i est grand, il faut s'interroger si c'est un point mal prédit ou un point atypique

$$|t_i| > t_{\alpha} \quad t_{\alpha} \approx 2-3$$

■ Résidu studentisé $t_i^* = \frac{\hat{\epsilon}_i}{\hat{\sigma}_{-i} \sqrt{1-h_{ii}}} = t_i \sqrt{\frac{n-1-p}{n-p-t_i^2}}$

Heuristique similaire

$$|t_i^*| > t_{\alpha} \quad t_{\alpha} \approx 2-3 \quad \text{"} \approx F(n-1-p) \text{"}$$

Rq Si le point est mal prédit ou atypique il a une influence forte sur $\hat{\sigma} \Rightarrow |t_i^*| > |t_i| \Rightarrow$ plus sensible

Rq: on a toujours $|\epsilon_i^*| \geq |\epsilon_i|$

■ Variabilité de la prédiction quand on retire i

Distance of fits

$$DEFFITS_i = \frac{\hat{y}_i - \hat{y}_{(i),i}}{\hat{\sigma}_{(i)} \sqrt{h_i}} = \epsilon_i^* \sqrt{\frac{h_i}{1-h_i}} \quad (\sim N(0,1))$$

Relation entre ϵ_i^* et h_i

Test lorsque $|DEFFITS_i| > c \sqrt{\frac{p}{n}} \quad c \approx 2-3$

■ Variabilité de la prédiction quand on retire i

Distance de Cook

$$D_i = \frac{\|\hat{y} - \hat{y}_{(i)}\|^2}{\hat{\sigma}^2 p} = \frac{\epsilon_i^2}{p} \frac{h_i}{1-h_i} \quad \left(\approx \frac{\chi^2(p)}{p} \right)$$

Test $D_i > 1$

ou $D_i > \frac{4}{n-p}$

■ Variabilité de l'estimation des coefficients quand on retire i

$$DEFBETAS_{j,i} = \frac{\hat{\beta}_j - \hat{\beta}_{(i),j}}{\hat{\sigma}_{(i)} \sqrt{(X^*X)^{-1}_{jj}}} \quad (\sim N(0,1))$$

$$= \epsilon_i^* \left[\frac{[(X^*X)^{-1} X^*]_{ji}}{\sqrt{(X^*X)^{-1}_{jj} (1-h_{ii})}} \right]$$

$$DEFBETAS_{j,i} > 1$$

$$DEFBETAS_{j,i} > \frac{c}{\sqrt{n}}$$

⇒ Mesure l'influence combinée par composante

■ Variabilité de la variance

$$V(\hat{\beta}) = \hat{\sigma}^2 \text{det}(X^*X)^{-1} \Rightarrow \text{règles du bon$$

$$\text{COVRATIO}_i = \frac{V(\hat{\beta}_{(i)})}{V(\hat{\beta})}$$

$$= \left[\frac{n-2-p}{n-p} + \frac{\epsilon_i^2}{n-p} \right] \frac{1}{(1-h_{ii})} \sqrt{\frac{p}{n-p}}$$

Comparer en rapport à 1

< 1 dégressif
> 1 croissant

$$\text{Test } |Cov(\hat{\beta}_0, -1)| > \epsilon^D \frac{p}{n}$$

$C \sim 2 \rightarrow$

3 / Résidus partiels,

outil de diagnostic

$$y_i = x_i^* \beta + \epsilon_i$$

$$\Rightarrow y_i = x_{i(i)}^* \beta_{(i)} + x_i \beta_i + \epsilon_i$$

$$\epsilon_{i(i)} = y_i - x_{i(i)}^* \beta_{(i)} = x_i \beta_i + \epsilon_i$$

$\Rightarrow$ on se ramène à 1 pb à ce variable

En pratique

$$\hat{\epsilon}_{i(i)} = y_i - x_{i(i)}^* \hat{\beta}_{(i)}$$

avec $\hat{\beta}_{(i)}$ estimé soit dans le modèle complet

soit dans le modèle sans x_i
en omettant x_i

$$\text{L'observation de } \hat{\epsilon}_{i(i)} = y_i - x_{i(i)}^* \hat{\beta}_{(i)}$$

en fonction de $x_{i(i)}$ permet de visualiser
le reste de dépendance en x_i

$\Rightarrow$ Estimation "visuelle" de l'information restante

VIII Sélection de Modèles / Sélection de Variables

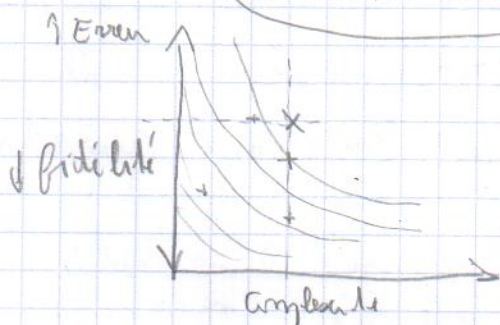
- Même question : comment choisir les régressions à utiliser ?

Objectifs $\neq$:
 prédiction $\rightarrow$ sélection de modèles
 interprétation $\rightarrow$ sélection de variables

- Principe de sélection = principe de parcimonie

Principe du rasoir d'Occam

Recherche d'un compromis entre fidélité aux observations et complexité.
 $\Rightarrow$ Biais / Variance



$\Rightarrow$ A même complexité, on choisit le modèle le plus fidèle

$\Rightarrow$ A même fidélité, on choisit le modèle le moins complexe

Pb Comment donner un ordre "total" sur $\mathbb{R}^2$?

- Ordre partiel : Test de Fisher

$$\mathcal{M}_0 \subset \mathcal{M}_1$$

$$\Rightarrow \frac{[RSS(0) - RSS(1)] / (p_1 - p_0)}{RSS(1) / (n - p_1)} > t_\alpha$$

$$RSS(0) + t_\alpha \frac{RSS(1)}{n - p_1} p_0 > RSS(1) + t_\alpha \frac{RSS(1)}{n - p_1} p_1$$

$\Rightarrow$ choix du modèle $\mathcal{M}_1$

Critère $C(p) = RSS(\mathcal{M}) + t_\alpha \hat{\sigma}^2 p$

$$\blacksquare R^2 = 1 - \frac{\|y - \hat{y}\|^2}{\|y - \bar{y}\|^2}$$

augmente avec la fidélité
 mais ne diminue pas avec la complexité

$\Rightarrow$ pas de compromis

$$R^2_q = 1 - \frac{\|Y - \hat{Y}\|^2 / (n-p)}{\|Y - \bar{Y}\|^2 / (n-1)}$$

prise en compte de la complexité avec le terme en $(n-p)$

$$R^2_a(\mathcal{M}_0) \leq R^2_a(\mathcal{M}_1)$$

$$\frac{\|Y - \hat{Y}_0\|^2}{n-p_0} \geq \frac{\|Y - \hat{Y}_1\|^2}{n-p_1}$$

$\hat{\sigma}_0^2 \geq \hat{\sigma}_1^2 \leadsto$ Variance dans le modèle prior
 $\Rightarrow$ Difficile à utiliser.

■ Principe d'optimisation on moyenne

qualité d'un estimateur mesurée par son erreur quadratique moyenne:

$$\begin{aligned} E(\|E(Y|X) - \hat{Y}\|^2) &= \underbrace{\|E(Y|X) - E(\hat{Y})\|^2}_{\text{biais}^2} + \underbrace{\text{Var}(\|\hat{Y}\|^2)}_{\text{Variance}} \\ &= \|E(Y) - P_{\text{Im}X} E(Y)\|^2 + \text{Var}(\|P_{\text{Im}X} Y\|^2) \\ &= \|E(Y) - P_{\text{Im}X} E(Y)\|^2 + \sigma^2 p \end{aligned}$$

pb: Terme de biais impossible à mesurer.

$\Rightarrow$ Estimation à l'aide de $\|Y - \hat{Y}\|^2$

$$\begin{aligned} E(\|Y - \hat{Y}\|^2) &= E(\|Y - P_{\text{Im}X} Y\|^2) \\ &= \|E(Y) - P_{\text{Im}X} E(Y)\|^2 + \text{Var}(\|Y - P_{\text{Im}X} Y\|^2) \\ &= \|E(Y) - P_{\text{Im}X} E(Y)\|^2 + \sigma^2 (n-p) \end{aligned}$$

$$\Rightarrow E(\|E(Y) - \hat{Y}\|^2) = E(\|Y - \hat{Y}\|^2) + (2p-n) \sigma^2$$

$\Rightarrow$ Utilisation du critère: $\|Y - \hat{Y}\|^2 + (2p-n) \sigma^2$

$\Rightarrow$ pb estimation de σ^2 .

$$\begin{aligned} \Rightarrow \text{choix } \hat{\sigma}^2 &= \frac{\|Y - \hat{Y}\|^2}{n-p} \Rightarrow \|Y - \hat{Y}\|^2 + (2p-n) \hat{\sigma}^2 = \frac{p}{n-p} \|Y - \hat{Y}\|^2 \\ &\quad \text{insuffisant en} \quad \text{Au lieu de si pas de biais} \quad = p R^2_a \end{aligned}$$

⇒ plus stricte que R^2_q

⇒ Forme classique : C_p de Mallows

$$C_p = \frac{\|Y - \hat{Y}\|^2}{\sigma^2} + 2p - m$$

Si l'estimateur est sans biais $C_p \geq p$

Règle classique : minimiser C_p

Règle du loir : choisir le modèle le plus simple tel que $C_p \leq p$.

Vraisemblance et généralisation

$$\ln(Y; (\beta, \sigma^2); \mathcal{M})$$

$$= -\frac{n}{2} \ln \sigma^2 - \frac{n}{2} \ln 2\pi - \frac{1}{2\sigma^2} \|Y - X\beta\|^2$$

Valeurs en le maximum de vraisemblance

$$\hat{\beta} = (X^T X)^{-1} X^T Y \quad \hat{\sigma}^2 = \frac{\|Y - X\hat{\beta}\|^2}{n} \neq \text{MCO}$$

$\ln(\text{ML}) :$

$$\ln(Y; (\hat{\beta}, \hat{\sigma}^2); \mathcal{M})$$

$$= -\frac{n}{2} \ln \hat{\sigma}^2 - \frac{n}{2} [\ln 2\pi + 1]$$

$$= -\frac{n}{2} \ln \left(\frac{\|Y - \hat{Y}\|^2}{n} \right) - \frac{n}{2} [\ln 2\pi + 1]$$

ne dépend que de l'attaché avec données (comme la R^2)

⇒ nécessité d'une pénalisation

choix classique : $-2\ln(\text{ML}) + 2\lambda_m |p|$

variable sur le choix de λ_m

$$\text{AIC} : \lambda_m = 2$$

$$\text{BIC} : \lambda_m = \log n$$

■ AIC : An Information Criterion
Akaike Information Criterion

Idée : Bon modèle est un modèle tel

quel $P(X; \hat{\theta}(y); \mathcal{M})$ soit proche de $P_0(X)$
en moyenne

⇒ Mesure par la divergence de Kullback-Leibler

$$KL(P_0, P(X; \hat{\theta}(y); \mathcal{M}))$$

$$= E_{X \sim P_0} \left[\ln \frac{P_0(X)}{P(X; \hat{\theta}(y); \mathcal{M})} \right]$$

$$= - E_{X \sim P_0} [-\ln P(X; \hat{\theta}(y); \mathcal{M})] + E_{X \sim P_0} [\ln P_0(X)]$$

$$E_{X \sim P_0} KL(P_0, P(X; \hat{\theta}(y); \mathcal{M}))$$

$$= - E_{X \sim P_0} E_{Y \sim P_0} [\ln P(X; \hat{\theta}(y); \mathcal{M})]$$

$$+ \underbrace{E_{X \sim P_0} [\ln P_0(X)]}_{cte}$$

Pb. Comment estimer $E_{Y \sim P_0} E_{X \sim P_0} [-\ln P(X; \hat{\theta}(y); \mathcal{M})]$

à partir de $-\ln P(Y; \hat{\theta}(y); \mathcal{M})$

Idée de la preuve

$$\theta_0 = \argmin E_{X \sim P_0} [-\ln P(X; \theta; \mathcal{M})]$$

$$E_{Y \sim P_0} E_{X \sim P_0} [-\ln P(X; \hat{\theta}(y); \mathcal{M})] - E_X [-\ln P(X; \theta_0; \mathcal{M})] \\ \approx d/2 \quad (d \text{ } \chi^2)$$

$$E_Y [-\ln P(Y; \theta_0; \mathcal{M})] - [-\ln P(Y; \theta_0(y); \mathcal{M})] \\ \approx d/2$$

⇒ χ^2 qui mesure...

$$E_{Y \sim P_0} E_{X \sim P_0} [-\ln P(X; \hat{\theta}(Y); M)]$$

$$\approx -\ln P(Y; \hat{\theta}(Y); M) + d = \text{AIC}(M) / 2$$

* Validation croisée. $\text{AIC}(M) = -2 \ln P(Y; \hat{\theta}(Y); M) + 2d$

BIC: Bayesian Information Criterion

Idée: Modélisation bayésienne et choix du modèle le plus probable étant donnée Y .

$$P(X; \theta; M) = P(Y | \theta; M) P(\theta | M) P(M)$$

$$= P(M, \theta | Y) P(Y)$$

$$P(M | Y) = \int P(M, \theta | Y) dP_\theta$$

$$= \int \frac{P(Y | \theta; M) P(\theta | M) P(M)}{P(Y)} P(\theta | M) d\theta$$

$$= \frac{P(M)}{P(Y)} \int P(Y | \theta; M) P^2(\theta | M) d\theta$$

Si $P(M) = \text{cte}$ il suffit de déterminer l'intégrale

$$\int P(Y | \theta; M) P^2(\theta | M) d\theta = \int e^{-n \left[\frac{1}{n} \left(-\ln P(Y | \theta; M) - 2 \ln P(\theta | M) \right) \right]} d\theta$$

lm: $\int e^{-n L(u)} du \approx e^{-n L(u^*)} \left(\frac{2\pi}{n} \right)^{\frac{d}{2}} | -L''(u^*) |^{-1/2}$

$$L = \frac{1}{n} [-\ln P(Y | \theta; M) - 2 \ln P(\theta | M)]$$

$$-L''(\theta^*) = + \frac{1}{n} \frac{\partial^2}{\partial \theta^2} [-\ln P(Y | \theta^*; M)] + \frac{2}{n} \frac{\partial^2}{\partial \theta^2} \ln P(\theta | M)$$

$$\hookrightarrow \text{GA}(\theta^*) = O\left(\frac{1}{n}\right)$$

$$- \ln P(\mathcal{M} | Y)$$

$$\approx - \ln P(Y | \mathcal{M}, \theta^*) - 2 \ln P(\theta^* | \mathcal{M})$$

$$+ \frac{d}{2} \log n - \frac{d}{2} \log 2\pi - \ln |A_n(\theta^*)|$$

$$\theta^* = \arg \max_{\theta} L(\theta) = \arg \max_{\theta} \left[-\frac{1}{n} \ln P(Y | \theta, \mathcal{M}) - \frac{2}{n} \ln P(\theta | \mathcal{M}) \right]$$

$$\rightarrow \hat{\theta} = \arg \max_{\theta} \left[\frac{1}{n} \ln P(Y | \theta, \mathcal{M}) \right]$$

$$\Rightarrow - \ln P(\mathcal{M} | Y) \approx - \ln P(Y | \mathcal{M}, \hat{\theta}) + \frac{d}{2} \log n$$

$$\underbrace{- 2 \ln P(\theta^* | \mathcal{M}) - \frac{d}{2} \log 2\pi - \frac{1}{2} \ln |A_n(\theta^*)|}_{O(1)}$$

$O(1)$

$$\Rightarrow \text{BIC}(\mathcal{M}) = - 2 \ln P(Y | \mathcal{M}, \hat{\theta}) + d \log n$$

$$\approx - \ln P(\mathcal{M} | Y)$$

Rq: ne dépend pas de $P(\theta | \mathcal{M})$

Prene en compte l'a priori sur les modèles

en ajoutant un terme $- \ln P(\mathcal{M})$

■ Comparaison

$$\begin{array}{ccc} R_a & (C_p) \text{ AIC} & \text{BIC} \\ \leftarrow & \text{complexité des modèles décroît} & \end{array}$$

Rq: $\text{AIC} \rightarrow C_p$ si σ^2 est connu.

lin avec F sur les modèles emboîtés.

■ Exploration exhaustive

→ Essayer tous les modèles
et choisir celui qui minimise le critère choisi.

1b : Algorithmique combinatoire → court exposé avec le mt de recherche

■ Exploration intelligente.

→ Explorer l'ensemble des modèles en ajoutant / enlevant des variables explicatives en fonction du critère.

→ Stratégie "gloutonne" : moins coûteuse mais difficile de garantir l'optimalité.

■ Partial least Square

on ajoute les régressions au fur et à mesure en fonction de leur corrélation avec le résidu partiel

on choisit ensuite des la famille explorée le meilleur

■ Modèle à pénalisation : si $X^T X$ n'est pas inversible

$$\|Y - X\beta\|^2 + \lambda \|\beta\|_p$$

Si $p=0$: critère de sélection de modèle

Si $p=2$: critère de régularisation de type Tychonov

$p=1$: sélection et algorithme efficace

⇒ LASSO sélection + résultat théorique

■ Validation croisée : Test des modèles avec une approche similaire à AZC en découpant les données en apprentissage et vérification.

Modèle linéaire généralisé

⇒ Généralisation du modèle linéaire

Famille de loi P_θ : gaussienne, binomiale, poisson

$$\eta_i = x_i^* \beta$$

$$\Rightarrow y_i \sim P_{\theta_i} \Rightarrow g(E(y_i)) = x_i^* \beta$$

↑
fonction de lien

ex P_θ = gaussien, $g(x) = x \Rightarrow$ Modèle linéaire gaussien

P_θ Bernoulli, $g(x) = \ln \frac{x}{1-x} \Rightarrow$ logit

P_θ Poisson, $g(x) = \ln x$

⇒ Famille "exponentielle", Maximum de vraisemblance, Tests

I Famille exponentielle

X_i variables indépendantes admettant des distributions issues d'une même structure exponentielle.

⇒ Densité admet une forme particulière

$$f(y_i; \theta_i, \phi) = e^{\underbrace{\frac{y_i \theta_i - \eta(\theta_i)}{\phi}}_{\text{dépend de } \theta_i} + \underbrace{w(y_i, \phi)}_{\text{ne dépend pas de } \theta_i}}$$

↑ ↑
variable variable
naturelle de nuisance

Rq $E(y_i; \theta, \phi) = \eta'(\theta)$

$$\text{Var}(y_i; \theta, \phi) = \phi \eta''(\theta)$$

$$\begin{aligned} \frac{\partial \ln f}{\partial \theta} &= \frac{y_i - \eta'(\theta)}{\phi} & E\left(\frac{\partial \ln f}{\partial \theta}\right) &= 0 \\ \frac{\partial^2 \ln f}{\partial \theta^2} &= -\frac{\eta''(\theta)}{\phi} & -E\left(\frac{\partial^2 \ln f}{\partial \theta^2}\right) &= E\left(\left(\frac{\partial \ln f}{\partial \theta}\right)^2\right) \end{aligned}$$

Rq L'estimateur du maximum de vraisemblance

de $X_1, \dots, X_n \sim f(\cdot; \theta; \phi)$ est $\hat{\theta}$ définie

$$\text{par } \eta'(\hat{\theta}) = \frac{1}{n} \sum X_i$$

⇒ Structure simple

⚠ estimateur de ϕ plus complexe...

Ex • Gaussienne

$$f(y; m, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2\sigma^2} (y-m)^2}$$
$$= e^{\left(\frac{ym - \frac{m^2}{2}}{\sigma^2} + \left[-\frac{y^2}{2\sigma^2} - \frac{1}{2} \ln 2\pi - \frac{1}{2} \ln \sigma^2\right]\right)}$$

$$\theta = m$$

$$\eta(m) = \frac{m^2}{2} \quad \phi = \sigma^2$$

$$\eta'(m) = m \Rightarrow E(y) = m$$

$$\eta''(m) = 1 \Rightarrow V(y) = \sigma^2$$

• Bernoulli

$$f(z, \pi) = \pi^z (1-\pi)^{1-z} = 1-\pi \left(\frac{\pi}{1-\pi}\right)^z$$
$$= e^{z \ln \frac{\pi}{1-\pi} + \ln(1-\pi)}$$

$$\theta = \ln \frac{\pi}{1-\pi} \Rightarrow e^\theta = \frac{\pi}{1-\pi} \Rightarrow (1-\pi)e^\theta = \pi$$

$$p=1 \Rightarrow \pi = \frac{e^\theta}{1+e^\theta}$$

$$\beta(z, \theta) = e^{z\theta - \ln(1+e^\theta)}$$

$$\eta(\theta) = -\ln(1+e^\theta) \quad (= \ln(1-\pi))$$

$$\eta'(\theta) = \frac{e^\theta}{1+e^\theta} \quad (= \pi = E(z)) \quad E(z) = \pi$$

$$\eta''(\theta) = \frac{e^\theta(1+e^\theta) - e^\theta e^\theta}{(1+e^\theta)^2} = \frac{e^\theta}{(1+e^\theta)^2} = \frac{e^\theta}{1+e^\theta} \cdot \frac{1}{1+e^\theta} \quad \left(= \pi \lambda (1-\pi) = V(z)\right)$$

• Poisson

$$f(z, \lambda) = \frac{\lambda^z e^{-\lambda}}{z!} = e^{z \ln \lambda - \lambda + \ln z!}$$

$$\theta = \ln \lambda \quad \phi = 1$$

$$\eta(\theta) = e^\theta \quad (= \lambda)$$

$$\eta'(\theta) = e^\theta \quad (= \lambda = E(z))$$

$$\eta''(\theta) = e^\theta \quad (= \lambda = E(z))$$

Poisson, Gamma ...

$$R_1 \quad v''(\theta) = \text{Var}(y) \Rightarrow v \text{ strictement convexe} \\ \Rightarrow E(y) = b'(\theta) \text{ est en ligne avec } \theta$$

Gaussien $E(y) = m = \theta$

Bernoulli $E(y) = \pi = \frac{e^\theta}{1+e^\theta} \Leftrightarrow \theta = \ln \frac{\pi}{1-\pi} = \ln \frac{E(y)}{1-E(y)}$ ↓ fonction canonique

Poisson $E(y) = \lambda = e^\theta \Leftrightarrow \theta = \ln \lambda = \ln E(y)$

II Modèles linéaires généralisés

- 1. Famille exponentielle pour y_i indépendants
- 2. Existence de régressum x_i et d'un coefficient β tel q.
- 3. $g(E(y_i|x_i)) = x_i^* \beta$ où g est une fonction croissante

Ex • Gaussienne : $g = \text{Id}$ (= fonction linéaire) $\Rightarrow$ modèle linéaire gaussien

• Bernoulli : $g = \ln \frac{x}{1-x}$ (= fonction croissante) $\Rightarrow$ modèle logit

$$g^{-1} = \frac{e^x}{1+e^x} \quad g^{-1}(1) = \frac{e^x}{1+e^x}$$

• $g = \Phi^{-1}(x)$ (= quantile gaussien) $\Rightarrow$ modèle probit

$$g^{-1}(x) = \Phi(x)$$

• $g(x) = \ln(-\ln(1-x))$ $\Rightarrow$ modèle log-log

$$g^{-1}(x) = 1 + e^{-e^{-x}}$$

• Poisson : $g(x) = \ln x$ $\Rightarrow$ modèle poissonien

$$g^{-1}(x) = e^x$$

avec lien logarithmique

$$\Rightarrow g(v'(\theta_i)) = x_i^* \beta \Leftrightarrow \theta_i = (v')^{-1}(g^{-1}(x_i^* \beta))$$

Ri si g fonction croissante $g(v'(\theta)) = \theta = x_i^* \beta$

$$\Rightarrow \theta_i = x_i^* \beta$$

□ Prédiction $\hat{\beta} \Rightarrow E(y) = g'(x^* \hat{\beta})$

□ Vraisemblance

$$l_n(y_i; \theta, \beta, \phi)$$

$$= \sum_{i=1}^n l_n(y_i; \theta(x_i, \beta), \phi)$$

$$= \sum_{i=1}^n \left[\frac{y_i \theta_i(\beta, x_i) - \eta(\theta_i(\beta, x_i))}{\phi} + w(y_i, \phi) \right]$$

$$\frac{\partial l_n}{\partial \beta_j} = \sum \frac{\partial l_n}{\partial \theta_i} \frac{\partial \theta_i}{\partial \beta_j}$$

$$= \sum \frac{y_i - \eta'(\theta_i(\beta, x_i))}{\phi} \frac{\partial \theta_i(\beta, x_i)}{\partial \beta_j}$$

$$\frac{\partial \theta_i}{\partial \beta_j} = \frac{1}{\eta''(\theta_i)} x_{ij} (g')'(x_i^* \beta)$$

$$\frac{\partial l_n}{\partial \beta_j} = \sum_{i=1}^n \frac{(y_i - \eta'(x_i^* \beta)) x_{ij}}{\phi \eta''(\eta^{-1}(\eta'(x_i^* \beta)))} (g')'(x_i^* \beta)$$

$\uparrow \eta''(\theta_i)$

$\Rightarrow$ Equations de maximum de vraisemblance

$$\frac{\partial l_n}{\partial \beta_j} = 0 \quad \forall j$$

$\Rightarrow$ Résolution unique possible

Rq: si $g = (\eta')^{-1}$ $(g')' = \eta''$

$$\Rightarrow \frac{\partial l_n}{\partial \beta_j} = \sum_{i=1}^n \frac{(y_i - \eta'(x_i^* \beta)) x_{ij}}{\phi}$$

$$\Rightarrow X^*(Y - \mu) = 0 \quad \text{soit} \quad \mu = \begin{pmatrix} \eta'(x_1^* \beta) \\ \eta'(x_n^* \beta) \end{pmatrix} = \begin{pmatrix} E(y_1) \\ \vdots \\ E(y_n) \end{pmatrix}$$

III Qualités d'ajustement et Tests

Qualités d'ajustement

• Dénna

$$D = -2(\mathcal{L} - \mathcal{L}_{\text{sat}}) \quad (\sim \text{RSS})$$

↑
vraisemblance
ou maximum de
vraisemblance
dans le modèle considéré
à p paramètres

↑
vraisemblance dans le modèle
solution à n paramètres

Asymptotiquement sous H_0 le modèle est vrai : $D \sim \chi^2(n-p)$

• Pearson

$$\text{Sous } H_0 \quad \chi^2 = \sum_{i=1}^n \frac{(y_i - g^{-1}(x_i^* \hat{\beta}))^2}{\phi v''((v')^{-1} g^{-1}(x_i^* \hat{\beta}))} \sim \chi^2(n-p)$$

Tests

• Rapports de vraisemblance

Pour deux modèles emboîtés $M_0 \subset M_1$

$$\text{Sous } H_0, \quad D_1 - D_0 = -2(\mathcal{L}_1 - \mathcal{L}_0) \sim \chi^2(p_1 - p_0)$$

si ϕ est connu

$$\sim F(p_1 - p_0, n-p)$$

Gaussian

• Test de Wald

$H_0: \mathbb{R}\beta = d$ avec $r = \text{nb de contraintes}$.

$$\Rightarrow \text{Sous } H_0 \quad (C\beta - d)^* \left(C (X^* W X)^{-1} C^* \right)^{-1} (C\beta - d) \sim \chi^2(r)$$

$$\text{ou } W \text{ est diagonale et vérifie } W_{ii} = \frac{1}{\phi v''((v')^{-1} g^{-1}(x_i^* \beta))} \left[(g')'(x_i^* \beta) \right]^2$$

Pb Estimer f à partir de l'observation
de $y_i = f(x_i) + \epsilon_i$

- on a un échantillon x_i $f(x_i) = \beta_0 + \beta_1 x_i$
ou plus généralement $f(x_i) = \beta_0 \sum_{k=0}^p \beta_k(x_i)$

Exemple $\beta_k(x_i) = x_i^k$

$\Rightarrow$ Approche paramétrique

Hyp $f \in \{f_\theta \text{ avec } \theta \in \mathbb{R}^p\} = \mathcal{M}$

$\Rightarrow \hat{f}_\theta$ avec $\hat{\theta} = \arg \min \sum |y_i - f(x_i)|^2$

- on a également étudié le cas de la famille

- Approche non paramétrique

Hyp $f \in \mathcal{C}^0 \Rightarrow$ Trop de solutions possibles!

$\approx$ on va chercher $\hat{f}_n$ dans un modèle $\mathcal{M}_n$
dont la complexité augmente avec $n \rightarrow \mathcal{C}^0$

- Ex $\mathcal{M}_n =$ polynôme de degré $k(n)$ avec $k(n) \rightarrow +\infty$

Sélection du degré par pénalisation

$L^2([0,1])$ ou validation croisée

$\mathcal{M}_n =$ polynôme trigonométrique	$\left\{ \begin{array}{l} \text{Biais} = \text{régularité} \\ \text{relat} \quad \text{grande} \\ \text{variance diminue} \end{array} \right.$
$\downarrow$ Même principe $L^2([0,1])$	

■ Binning = découpage en boîte

⇒ utiliser un modèle linéaire simple dans chaque boîte

⇒ constant / polynomial / ...

⇒ difficulté toute de boîte : biais $\rightarrow$ variance.

■ Pb : bord des boîtes $\Rightarrow$ discontinuité

⇒ solution de la régression locale

⇒ régression avec des poids voisins

⇒ Pb de discontinuité entre une entrée et la suite des points

⇒ Noyau à support compact continu / seules sont

les régressions locales

⇒ Pb d'ordre de la largeur de boîte Biais
du paramètre h Variance

■ Boîte qui se "peut" : condition de parcourir

⇒ spline

⇒ choix des moments $\Rightarrow$ sélection de modèle.

⇒ Problématique biais variance

(Formulation variationnelle?)

■ Ondelettes : ex de la box de Fourier.

"Pb de la boîte des boîtes"

$\Rightarrow$ modèles emboîtés + Box errant sur $L^2(\mathbb{R})$