

- Objectifs séance:
1. Récap cours et quiz
 2. Étude approfondie du MLE
 3. Étude des modèles exponentiels → PC3.

Modèle $(\mathcal{Z}, \mathcal{X}, \{P_\theta, \theta \in \Theta\})$

Estimateur: fct mes à valeurs de $\mathbb{R}^q$
 si $g(\theta) \in \mathbb{R}^q$.

EMM:

$$\forall \theta \in \Theta \quad \mathbb{E}_\theta(T_j(X)) = e_j(\theta)$$

$$\rightarrow \text{EMM}(\tau) / e_j(\hat{\theta}_n) = \frac{1}{n} \sum_{i=1}^n T_j(X_i)$$

Théorème: LCN / TCL

$\mathcal{Z}$ estimateur: le modèle est un $\mathcal{Z}$ estimat
 par un EMM.

EMV modèle à densité

$$\{g_\theta^{\otimes n}, \theta \in \Theta\}$$

$$\begin{aligned} \mathcal{L}_n(\theta, X) &= g_\theta^{\otimes n}(X) \\ &= \prod_{i=1}^n g_\theta(X_i) \end{aligned}$$

$$\begin{aligned} \underline{\text{EMV}} \quad \hat{\theta}_{\text{EMV}} &= \underset{\theta \in \Theta}{\operatorname{argmax}} \mathcal{L}_n(\theta, X) \\ &= \underset{\theta \in \Theta}{\operatorname{argmax}} \mathcal{L}_n(\theta, X) \end{aligned}$$

s'ils existent.

Bonne EMV $\subset \Pi$ -estimateurs.

$$\frac{1}{n} \sum_{i=1}^n \log g_\theta(X_i) \xrightarrow{P_{\theta^*}} \mathbb{E}_{\theta^*}(\log g_\theta(X))$$

$$\begin{aligned} \theta^* &= \underset{\theta}{\operatorname{argmax}} \underbrace{\Pi(\theta, \theta^*)}_{d_{\theta^*} = KL(P_{\theta^*}, P_\theta)} \\ &= \underset{\theta}{\operatorname{argmax}} \Pi(\theta, \theta^*) \end{aligned}$$

Pl de vue Π estimat: $fct \Pi(\theta, \theta^*, X) = \mathbb{E}_{\theta^*}(\log g_\theta(X))$

$$\hat{\theta}_{\text{EMV}} = \underset{\theta}{\operatorname{argmax}} \frac{1}{n} \sum_{i=1}^n \log g_\theta(X_i)$$

Quizz:

* Soit $A \in \mathbb{R}^{n \times d}$ et $X \sim \mathcal{N}(\mu, \Sigma)$ $\mu \in \mathbb{R}^d$, $\Sigma \in \mathbb{R}^{d \times d}$

Loi de $AX \sim \mathcal{N}(A\mu, A\Sigma A^T)$

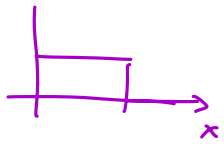
* Donner un exemple de modèle où $EMV \neq EMM(T)$

$(\mathbb{R}^n, -, \{\mathcal{U}[0, \theta]^{\otimes n}, \theta > 0\})$

$EMV : \max_{i=1..n} X_i = \hat{\theta}_{EMV}$

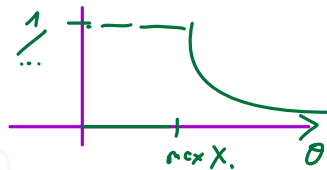
$EMM(T=Id) : \hat{\theta}_n = 2\bar{X}_n$

$$p_\theta(x) = \frac{1}{\theta} \mathbb{1}_{0 \leq x \leq \theta}$$



Preuve : $L_n(\theta, X) = \prod_{i=1}^n p_\theta(X_i) = \prod_{i=1}^n \frac{1}{\theta} \mathbb{1}_{X_i \leq \theta}$

$\hat{\theta}_{EMV} = \arg\max L_n(\theta, X) = \max(X_i) = X_{(n)} = \frac{1}{\theta^n} \mathbb{1}_{\max_{i=1..n} X_i \leq \theta}$



C'est l'objet de l'exo 3 de montrer que pour une large classe de modèle (les modèles exponentiels), l'EMV coïncide à l'EMM pour un T^* naturel.

* On considère le modèle de Poisson : $\{\mathbb{N}, \mathcal{P}(\mathbb{N}), \{P_\lambda, \lambda > 0\}\}$

Proposer plusieurs EMM.

$$P_\lambda(X=k) = \frac{\lambda^k}{k!} e^{-\lambda}$$

* Soit $X \in \mathbb{R}^{n \times d}$, $X = \begin{pmatrix} x_1 & x_2 & \dots & x_d \\ \vdots & \vdots & & \vdots \end{pmatrix} \in \mathbb{R}^{n \times d}$ $n \geq d$.

→ Exprimer $\text{Im } X$. Majorer $\text{rg}(X)$.

→ $\text{Ker } X = (\text{Im}(X^T))^{\perp}$. Montrer que $(x_1, \dots, x_d)$ famille libre $\Leftrightarrow X^T X$ inversible.

→ Montrer alors que $\forall z; \hat{f}_{\text{Im}(X)}^{\perp}(z) = X(X^T X)^{-1} X^T z$

→ $\text{Im } X = \{X\alpha, \alpha \in \mathbb{R}^d\} = \left\{ \sum_{i=1}^d \alpha_i x_i, \alpha \in \mathbb{R}^d \right\} = \text{Vect}(x_1, \dots, x_d)$
 $\text{rg}(X) = \min(n, d) = d$ $n \geq d$.

→ $\text{Ker}(X) = \text{Im}(X^T)^{\perp}$. $(X\alpha)^T \beta = \alpha^T (X^T \beta)$
 $\parallel \parallel$

$\alpha \in \text{Ker}(X) \Leftrightarrow X\alpha = 0 \Leftrightarrow \forall \beta, \langle X\alpha, \beta \rangle = 0 \Leftrightarrow \forall \beta, \langle \alpha, X^T \beta \rangle = 0 \Leftrightarrow \alpha \in (\text{Im } X^T)^{\perp}$

→ $(x_1, \dots, x_d)$ famille libre $\Leftrightarrow X^T X$ inversible.

$(x_1, \dots, x_d)$ f.o. libre $\Leftrightarrow \text{rg}(X) = d \Leftrightarrow \text{Ker}(X) = \{0\} \Leftrightarrow X \text{ inj} \Leftarrow X^T X \text{ bij}$

$X \text{ inj} \Rightarrow (X^T X) \text{ inv.} : u \in \text{Ker}(X^T X) : X^T X u = 0 \Rightarrow u^T X^T X u = 0$
 $\Rightarrow \|Xu\|^2 = 0 \Rightarrow Xu = 0$

D'où $\text{Ker}(X^T X) = \{0\}$

$\Rightarrow u = 0$
 $X \text{ inj}$

→ $P_{\text{Im}(X)}^+(y)$ est caractérisé par: $P_{\text{Im}(X)}(y) \in \text{Im}(X)$ $(*)$
 $y - P_{\text{Im}(X)}(y) \in (\text{Im}(X))^{\perp} = \text{Ker}(X^T)$

formule donnée: $P_{\text{Im}(X)}(y) = X(X^T X)^{-1} X^T y$
 $\in \text{Im}(X)$

$X^T(y - X(X^T X)^{-1} X^T y) \stackrel{(*)}{=} 0$

pr de retour

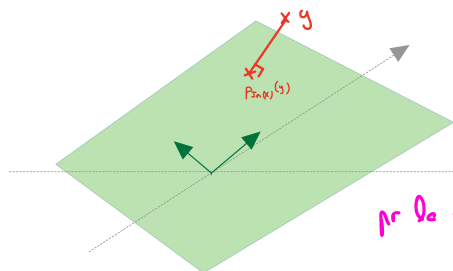
$P_{\text{Im}(X)}(y) = Xz$

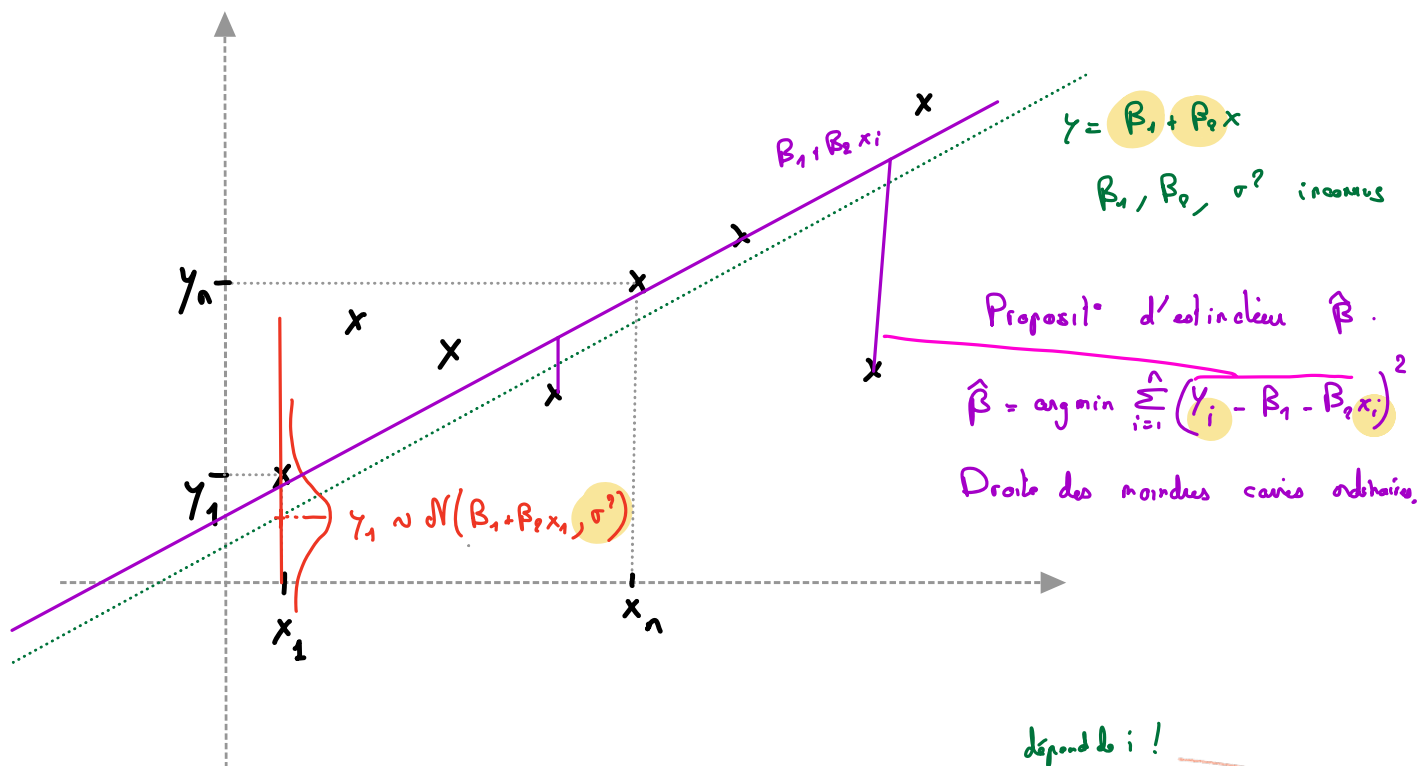
avec $z / (y - Xz) \in (\text{Im } X)^{\perp}$

$\Rightarrow \forall u, \langle y - Xz, Xu \rangle = 0 \Rightarrow (X^T y - X^T X z) = 0$

$\Rightarrow z = (X^T X)^{-1} X^T y$

$\Rightarrow P_{\text{Im}(X)}(y) = X(X^T X)^{-1} X^T y$





On considère le modèle $R^1, \mathcal{B}(R^1), \{P_\theta, \mathcal{L}^{\otimes n} = \bigotimes_{i=1}^n \mathcal{N}(\overbrace{B_1 + B_2 x_i}^{\text{dépend de } i!}, \sigma^2)\}$;
 Δ x_i ne s't pas les obs, ils s't fixes et connus.
 $\theta := (B_1, B_2, \sigma^2) \in \mathbb{R}^2 \times \mathbb{R}_+^*$

Les obs s't notées $(y_1 \dots y_n)^T = \mathbf{y} \rightarrow (y_i) \bigotimes_{i=1}^n \rightarrow$ "indep"

$\rightarrow$ non identiq distribués.
 $y_i \sim \mathcal{N}(B_1 + B_2 x_i, \sigma^2)$

ⓓ d'Exo : On va mq la droite des moindres carrés :

(qu'on ome pcq "pendre les g'des valeurs", "différ. et résul exploit")

on exactent l'EMV de le modèle gaussien homoscédastiq

σ^2 ne dép. pas de i .

The concept of regression comes from [genetics](#) and was popularized by Sir Francis Galton during the late 19th century with the publication of *Regression towards mediocrity in hereditary stature*.^[12] Galton observed that extreme characteristics (e.g., height) in parents are not passed on completely to their offspring. Rather, the characteristics in the offspring regress towards a *mediocre* point (a point which has since been identified as the mean). By measuring the heights of hundreds of people, he was able to quantify regression to the mean, and estimate the size of the effect. Galton wrote that, "the average regression of the offspring is a constant fraction of their respective [mid-parental](#) deviations". This means that the difference between a child and its parents for some characteristic is proportional to its parents' deviation from typical people in the population. If its parents are each two inches taller than the averages for men and women, then, on average, the offspring will be shorter than its parents by some factor (which, today, we would call one minus the [regression coefficient](#)) times two inches. For height, Galton estimated this coefficient to be about 2/3: the height of an individual will measure around a midpoint that is two thirds of the parents' deviation from the population average.

Galton coined the term "regression" to describe an observable fact in the inheritance of multi-factorial [quantitative genetic](#) traits: namely that the offspring of parents who lie at the tails of the distribution will tend to lie closer to the centre, the mean, of the distribution. He quantified this trend, and in doing so invented [linear regression](#) analysis, thus laying the groundwork for much of modern statistical modelling. Since then, the term "regression" has taken on a variety of meanings, and it may be used by modern statisticians to describe phenomena of [sampling bias](#) which have little to do with Galton's original observations in the field of genetics.

Wiki: À l'instar de Galton auquel le lie une indéfectible amitié, Pearson pense qu'il est possible et même éminemment souhaitable d'améliorer la race humaine en sélectionnant et favorisant les plus doués de ses représentants comme le fait la sélection naturelle pour les animaux.

$$\left(\beta_1, \beta_2 x_i \right)_{i=1}^n$$

L'eugénisme n'est pas une théorie scientifique, en relevant plutôt de l'idéologie scientifique[16]. Cependant en ayant recours à des méthodes statistiques pour évaluer les problèmes de son époque (dégénérescence et transmission de caractères mesurables dans une population), Karl Pearson et Walter Weldon fondent la biométrie. Les eugénistes anglais ont ainsi conçu et amélioré des outils mathématiques qui seront largement repris par les généticiens[14].

$$Y = (y_1 \dots y_n), \quad X = \begin{pmatrix} 1 & x_1 \\ \vdots & \vdots \\ 1 & x_n \end{pmatrix}, \quad \beta = \begin{pmatrix} \beta_1 \\ \beta_2 \end{pmatrix} \Rightarrow Y \sim \mathcal{N}(X\beta, \sigma^2 I_n)$$

gauss indép

modèle équivalent

1) $M_q(X^T X)$ est bijective

$$(X^T X) \text{ bij} \xRightarrow{\text{repl}} X \text{ inj} \iff \begin{pmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{pmatrix}, \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \text{ libre} \iff (x_i)_{i=1}^n \text{ non tous égaux!}$$

Rq : si X est non inject. : q se passe-t-il.

X matrice des var. explicatives : si les colonnes s't lin dép. je ne peux pas identifier $\beta_1 \dots \beta_p$.

Modèle est non identifiable : $\theta \longrightarrow \mathcal{N}(X\beta, \sigma^2 I)$ non inj
si $\beta \rightarrow X\beta$ non inj !

2) Déterminer l'EMV :

$$\text{sd 1} \quad L_n(\theta, Y) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi} \sigma} e^{-\frac{(y_i - \beta_1 - \beta_2 x_i)^2}{2\sigma^2}}$$

$$L_n(\theta, Y) = \frac{1}{\sqrt{(2\pi)^n \sigma^{2n}}} \exp - \frac{\sum_{i=1}^n (y_i - (\beta_1 + \beta_2 x_i))^2}{2\sigma^2} \quad \text{densité de } y_i \text{ / l.c.b.}$$

$l_n(\theta, Y) = \dots \rightarrow$ on différencie en $\beta_1, \beta_2, \sigma^2$.

Sol 2 = sd pro

densité de Y .

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

$$L_n(\theta, Y) = \frac{1}{\sqrt{(2\pi)^n \det \Sigma}} \exp - \frac{(Y - X\beta)^T \Sigma^{-1} (Y - X\beta)}{2}$$

(c'est pareil ou similaire)

$$\exp \left(- \frac{\|Y - X\beta\|^2}{2\sigma^2} \right)$$

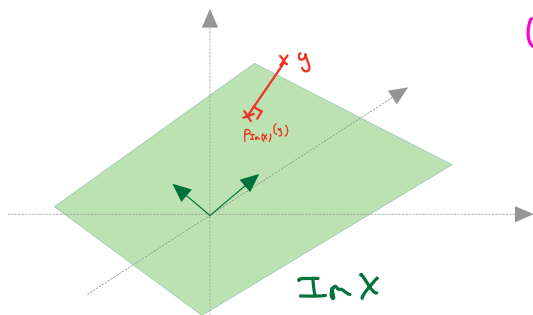
Il donne

$$l_n(\theta, y) = -\log(\ell\pi)^{\frac{n}{2}} - \frac{n}{2} \log(\sigma^2) - \frac{\|y - X\beta\|^2}{2\sigma^2}$$

Que dire de $\arg\max_{\theta=(\beta, \sigma^2)} l_n(\theta, y)$.

Rq: β correspond à la droite des moindres carrés!!
 $\sum_{i=1}^n (y_i - (\beta_0 + \beta_1 x_i))^2$

2/ $\arg\min_{\beta} \|y - X\beta\|^2$ ne dépend de la valeur de σ :
 Idée: il faut en rajouter
 Idée: ne pas différencier trop vite



(1) $\Rightarrow$ $X\hat{\beta}_{ENV} = P_{I_n(X)}(y) = X(X^T X)^{-1} X^T y$
 $\Rightarrow \hat{\beta}_{ENV} = (X^T X)^{-1} X^T y$

$X\beta$ parcourt $I_n(X)$

$\arg\min_{\beta \in I_n(X)} \|y - z\|^2 = P_{I_n(X)}(y)$

3) Loi de $\hat{\beta}_{ENV}$?

$\underline{\text{Sous } P_0}; y \sim N(X\beta, \sigma^2 I)$

Loi
 $z \sim N(\mu, \Sigma)$
 $Az \sim N(A\mu, A\Sigma A^T)$

$\hat{\beta}_{ENV} \sim N(\underbrace{(X^T X)^{-1} X^T}_{\beta} X\beta; \sigma^2 \cancel{(X^T X)^{-1} X^T X (X^T X)^{-1}})$

$\underline{\text{Sous } P_0} \Rightarrow \hat{\beta}_{ENV} \sim N(\beta; \sigma^2 (X^T X)^{-1})$

$\forall \theta; E_{\theta}(\hat{\beta}_{ENV}) = \beta$: c'est un est. sans biais!!

4) $\pi_1 \hat{\beta}_{ENV,1} \perp \hat{\beta}_{ENV,2} \Rightarrow \sum_{i=1}^n x_i = 0$

$\hat{\beta}_{ENV,1} \perp \hat{\beta}_{ENV,2} \Leftrightarrow (X^T X)^{-1} \text{ diag} \Leftrightarrow (X^T X) \text{ diag} \Leftrightarrow \sum_{i=1}^n x_i = 0$
 $\begin{pmatrix} n & \sum x_i \\ \sum x_i & \sum x_i^2 \end{pmatrix}$

5) $\hat{\sigma}_{ENV}^2 = \arg\max_{\sigma^2} \underbrace{-\frac{n \log \sigma^2}{2} - \frac{\|y - X\hat{\beta}_{ENV}\|^2}{2\sigma^2}}_{g(\sigma^2)}$

$$\frac{\partial}{\partial \sigma^2} g(\sigma^2) = -\frac{n}{2\sigma^2} + \frac{\|Y - X\hat{\beta}\|^2}{2\sigma^4} = 0$$

$$\Leftrightarrow \hat{\sigma}_{\text{ML}}^2 = \frac{\|Y - X\hat{\beta}_{\text{ML}}\|^2}{n} = \frac{\|Y - P_{\text{span}(X)}(Y)\|^2}{n} = \frac{\|P_{\text{span}(X)^\perp}(Y)\|^2}{n}$$

2 questions : loi de $\hat{\sigma}_{\text{ML}}^2$; $\hat{\sigma}_{\text{ML}}^2 \perp \hat{\beta}_{\text{ML}}$

Théorème : Théorème de Cochran

Soit X un vecteur aléatoire gaussien de $\mathbb{R}^n$ de loi $\mathcal{N}(0_{\mathbb{R}^n}, \sigma^2 \text{Id}_n)$. Soit F un sev de $\mathbb{R}^n$ de dimension d , $F^\perp$ son orthogonal et $P_F, P_{F^\perp}$ les matrices des projections orthogonales sur F et $F^\perp$. Alors :

- les vecteurs aléatoires $P_F X, P_{F^\perp} X$ sont indépendants et de lois respectives $\mathcal{N}(0_{\mathbb{R}^n}, \sigma^2 P_F), \mathcal{N}(0_{\mathbb{R}^n}, \sigma^2 P_{F^\perp})$;
- les variables aléatoires réelles $\|P_F X\|^2, \|P_{F^\perp} X\|^2$ sont indépendantes et de lois respectives $\sigma^2 \chi^2(d), \sigma^2 \chi^2(n-d)$.

Preuve $Z \sim \mathcal{N}(0, \sigma^2 I_n)$

Cas n°1 $F = \text{Vect}(e_1, \dots, e_d)$; $P_F(z_1, \dots, z_n) = (z_1, \dots, z_d, 0, \dots, 0)$

$P_{F^\perp}(\text{---}) = (0, \dots, 0, z_{d+1}, \dots, z_n)$

$(z_1, \dots, z_n)$ iid. $\Rightarrow P_F(Z) \perp P_{F^\perp}(Z)$!

$$\|P_F(Z)\|^2 = \sum_{i=1}^d z_i^2 \sim \sigma^2 \chi^2(d)$$

↑
def.

Cas général je considère une BON adaptée à F et $F^\perp$ et U la matrice de chgt de base.

soit $W = UZ$. $\rightarrow P_F(Z) = (w_1, \dots, w_d, 0, \dots, 0)$

$$W \sim \mathcal{N}(0, \sigma^2 U^T U) = \mathcal{N}(0, \sigma^2 I_n).$$

On se ramène au cas $n \geq 1$.

$$\text{On a } X \hat{\beta}_{\text{ENV}} = P_{\text{Im}(X)}(Y) \quad \downarrow \quad \hat{\sigma}_{\text{ENV}}^2 = \frac{\|P_{\text{Im}(X)}^\perp(Y)\|^2}{n}$$

$$\Rightarrow \hat{\beta}_{\text{ENV}} \perp \hat{\sigma}_{\text{ENV}}^2$$

$\uparrow$ lemme de corollaire.

$$\text{Loi de } \hat{\sigma}_{\text{ENV}}^2 = \frac{\|P_{\text{Im}(X)}^\perp(Y)\|^2}{n} = \frac{\|P_{\text{Im}(X)}^\perp(Y - XB)\|^2}{n} \sim \sigma^2 \chi^2(?)$$

$$\Delta \quad Y \sim \mathcal{N}(XB; \sigma^2 I)$$

$\uparrow$ non centrée

$$\dim(\text{Im } X^\perp) = n - 2$$

$$\hat{\sigma}_{\text{ENV}}^2 \sim \frac{\sigma^2 \chi^2(n-2)}{n}$$

Gosset : cas particuliers avec 1 seule v.ar. explicative

$$X = \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix} \quad \beta = \mu$$

$$Y \sim \mathcal{N}(\mu \vec{1}, \sigma^2 \text{Id}) \quad \rightarrow \quad \hat{\mu}_n = \bar{Y}_n = \underbrace{(\vec{1}^T \vec{1})^{-1}}_{\frac{1}{n}} \underbrace{\vec{1}^T Y}_{\sum Y_i}$$

$$\hat{\sigma}_{\text{ENV}}^2 = \frac{1}{n} \sum (Y_i - \bar{Y}_n)^2$$

On retrouve l'exo 2 PC 1 !