

Apprentissage supervisé

→ 1^{er} modèle : régression linéaire

$$h_{\beta}(x) = \beta_0 + \beta_1 x_1 + \dots + \beta_D x_D$$

β_j : coefficients de régression

Ensemble de données d'entraînement $\mathcal{D} = \{x^{(i)}, t^{(i)}\}_{i=1}^N$

Fonction de coût $\mathcal{L}(\beta) = \frac{1}{N} \sum_{i=1}^N (h_{\beta}(x^{(i)}) - t^{(i)})^2$

→ Apprentissage : minimisation de $\mathcal{L}(\beta)$

$$\vec{\beta} \in \mathbb{R}^{D+1}$$

$$\vec{x} \in \mathbb{R}^D$$

$$y(x) \in \mathbb{R}$$

$$\tilde{x} = [1, \vec{x}] \in \mathbb{R}^{D+1}$$

→ 2 approches : Descent du gradient

$$\beta_{\text{init}} : \beta \leftarrow \beta - \eta g_{\beta}^{\text{grad}} l(\beta)$$

Résolution des équations normales

$$\rightarrow (\tilde{X}^T \tilde{X}) \vec{\beta} = \tilde{X}^T \vec{t} \Rightarrow \vec{\beta}_{OLS} = (\tilde{X}^T \tilde{X})^{-1} \tilde{X}^T \vec{t}$$

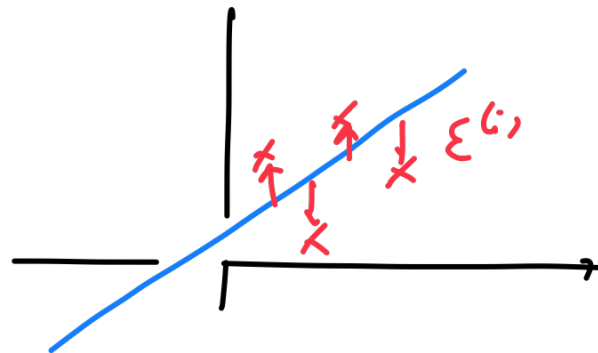
$$\begin{aligned} \underline{\tilde{X}} &= \begin{bmatrix} 1 & x_1^{(2)} & x_2^{(2)} & \dots & x_D^{(2)} \\ \vdots & \vdots & \vdots & & \vdots \\ 1 & x_1^{(N)} & \dots & & x_D^{(N)} \end{bmatrix} \rightarrow \begin{bmatrix} x_1 & x_2 \\ x_1' & x_2' \end{bmatrix} \\ &= \begin{bmatrix} x_1^2 + (x_1')^2 & x_1 x_2 + x_1' x_2' \\ x_1 x_2 + x_1' x_2' & x_2^2 + (x_2')^2 \end{bmatrix} = \begin{bmatrix} x_1 & x_1' \\ x_2 & x_2' \end{bmatrix} \begin{bmatrix} x_1 & x_2 \\ x_1' & x_2' \end{bmatrix} \end{aligned}$$

→ Aujourd'hui

→ Motivations statistiques

→ Cas de $X^T X$ non inversible

(présence de dépendances au niveau des caractéristiques)



Motivation Statistique

Hypothèse #1 $y^{(i)} = \beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D + \epsilon^{(i)}$ $\epsilon^{(i)} \sim \mathcal{N}(0, \sigma^2)$

Hypothèse #2 $\epsilon^{(i)}$ i.i.d.
(indépendants et
identiquement distribués)

$$\mathcal{N}(0, \sigma^2) \sim \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{x^2}{2\sigma^2}\right)$$

$$t^{(i)} \sim \mathcal{N}(\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}, \sigma^2)$$

$$\sim \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2}{2\sigma^2}\right)$$

$$p(\{t^{(i)}\}_{i=1}^N | \{x^{(i)}\}_{i=1}^N, \beta) \sim \prod_{i=1}^N \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(t^{(i)} - h_{\beta}(x^{(i)}))^2}{2\sigma^2}\right)$$

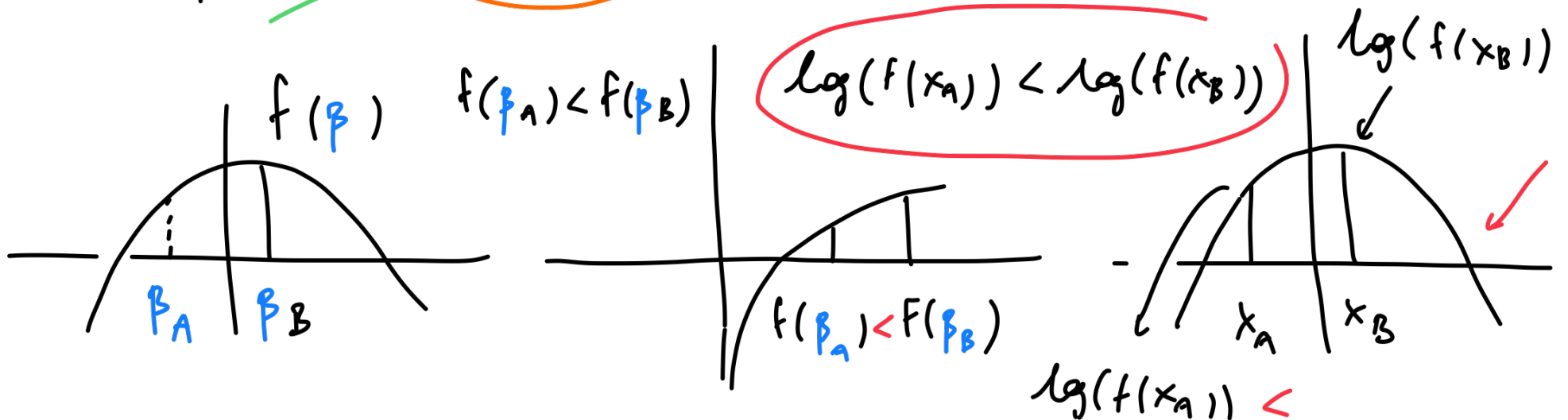
$$\rightarrow \hat{\beta}_{MLE} \rightarrow \operatorname{argmax}_{\beta} \prod_{i=1}^N \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2}{2\sigma^2}\right)$$

$\hat{\beta}_{MLE} \rightarrow$ estimateur de Maximum de vraisemblance

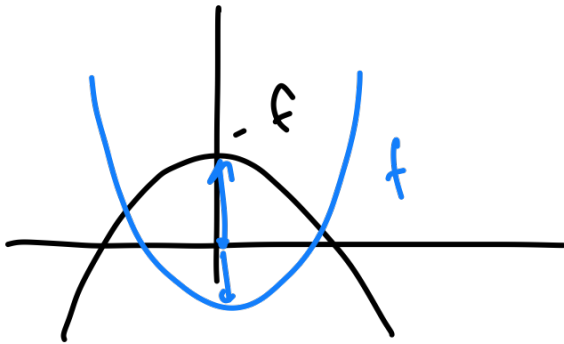
$$\hat{\beta}_{MLE} = \operatorname{argmax} \log \left(\prod_{i=1}^N \frac{1}{\sqrt{2\pi}\sigma} \exp \left(- \frac{(t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2}{2\sigma^2} \right) \right)$$

$$= \operatorname{argmax} \sum_{i=1}^N \left[\log \left(\frac{1}{\sqrt{2\pi}\sigma} \right) - \frac{(t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2}{2\sigma^2} \right]$$

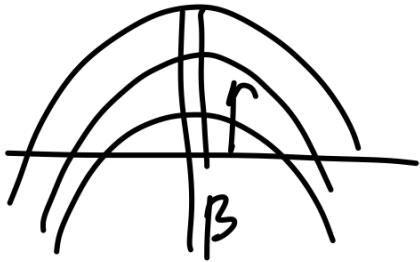
$$= \operatorname{argmax}_{\beta} \left(N \log \left(\frac{1}{\sqrt{2\pi}\sigma} \right) - \frac{1}{2\sigma^2} \sum_{i=1}^N (t^{(i)} - h_{\beta}(x^{(i)}))^2 \right)$$



$$\beta_{HLE} = \argmax_{\beta} \underbrace{- \frac{1}{2\sigma^2} \sum_{i=1}^N (t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2}_{\text{loss}}$$



$$\beta_{HLE} = \argmin_{\beta} \frac{1}{2\sigma^2} \sum_{i=1}^N (t^{(i)} - h_{\beta}(x^{(i)}))^2$$



$$\argmax_{\beta} f(\beta) + c$$

$$\argmax_{\beta} f(\beta)$$

Cas du $\tilde{X}^T \tilde{X}$ non inversible

$$\tilde{X} = \begin{bmatrix} 1 & x_1^{(2)} & \dots & x_D^{(2)} \\ \vdots & \vdots & & \vdots \\ 1 & x_1^{(N)} & \dots & x_D^{(N)} \end{bmatrix}$$

Traiter la dépendance

Portabilité #1 : Orthogonalisation successive (a.k.a Gram Schmidt)

→ Commencer avec $z_0^{(0)}$ première colonne de $\tilde{X}$

For $j = 1, \dots, D$

Calculer les coefficients $\hat{\gamma}_{lj} = \frac{\langle z_l, c_j \rangle}{\langle z_l, z_l \rangle}$

for $l = 0, \dots, j-1$

→ Et définir le nouveau z_j comme

$$z_j = c_j - \sum_{k=0}^{j-2} \hat{\gamma}_{kj} z_k$$

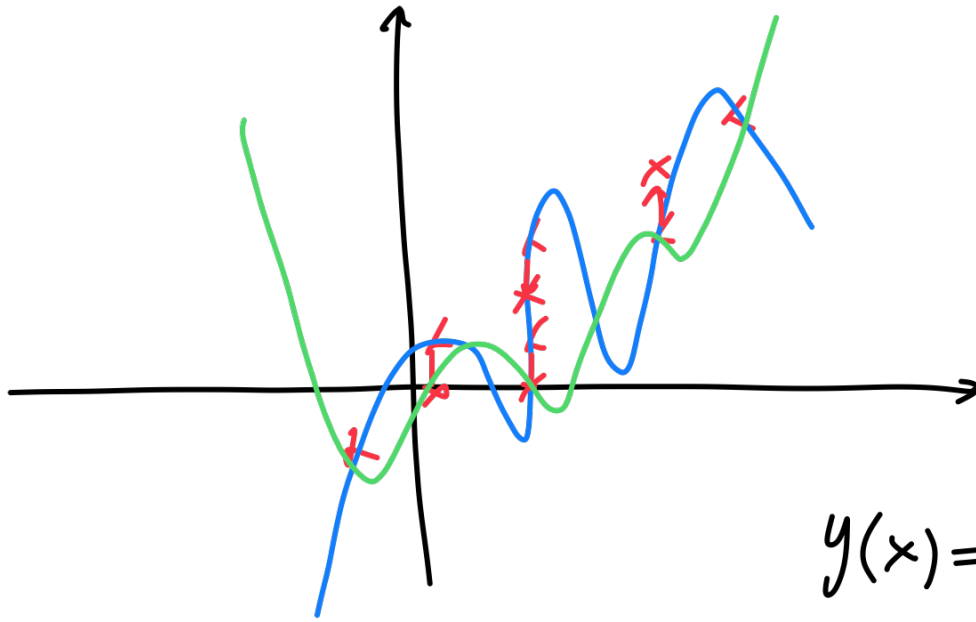
$$\rightarrow \underline{z_{D+2}} = c_{D+1} - \sum_{k=0}^D \hat{\gamma}_{kj} z_k$$

$\vec{z}_{D+2}$ est la seule colonne représentant c_{D+1}

$$\beta_{D+1} = \frac{\langle z_{D+2}, \vec{t} \rangle}{\underbrace{\langle z_{D+1}, z_{D+2} \rangle}} \rightarrow \text{si } z_{D+2} \text{ est petit}$$

(caractéristiques dépendantes)

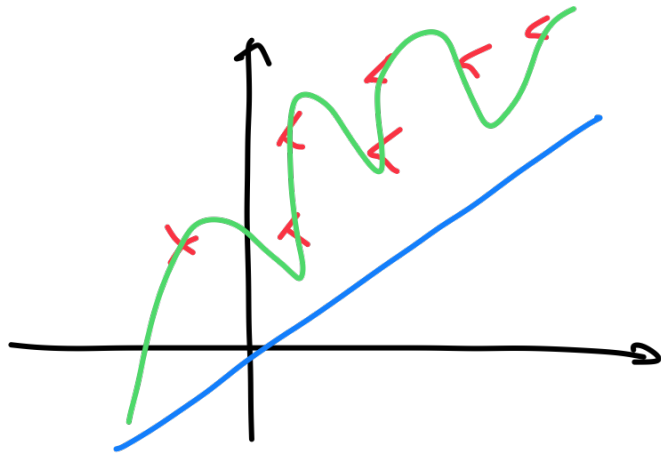
$\beta_{D+1} \gg \gg \Rightarrow$ estimation instable



$$y(x) = \beta_0 + \beta_1 x + \beta_2 x^2 + \dots +$$

$$\underbrace{\beta}_{OLS} = \underbrace{(\tilde{X}^T \tilde{X})^{-1}}_{= \quad =} \underbrace{\tilde{X}^T \vec{t}}_{= \quad =}$$

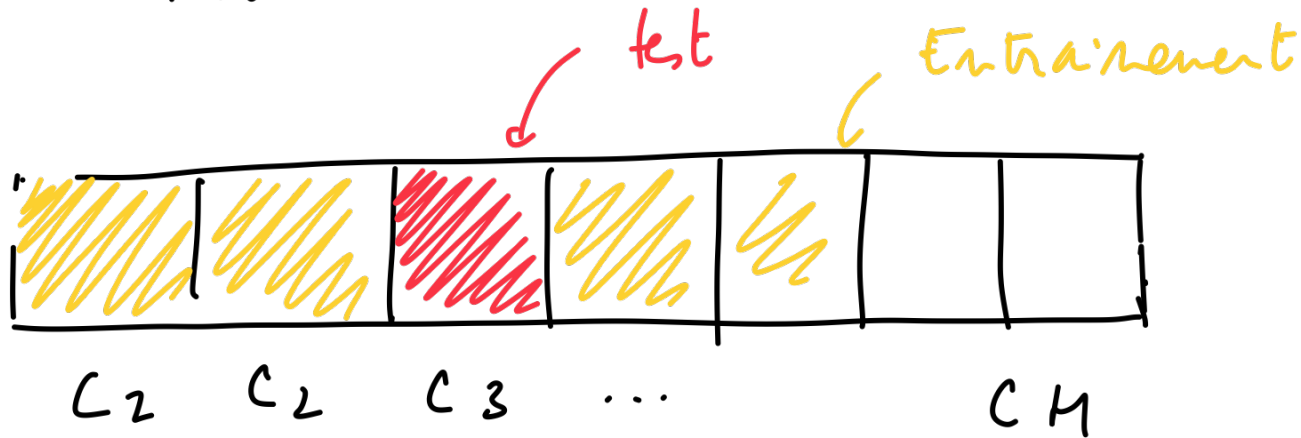
→ Alternative #2 : Sélection du meilleur sous ensemble



Nombre de sous-ensembles de taille k
choisis parmi un ensemble de taille
 D

$$\binom{D}{k}$$

→ Evaluation des $\binom{D}{k}$ modèles de taille k par
validation croisée



M Compartiments

$C_1, \dots, C_M$

$$|C_i| = \frac{N}{M}$$

Etape #2

$$\rightarrow \ell_{CV}(\beta) = \frac{1}{N} \sum_{i=1}^N \left(t^{(i)} - h_{\beta}^{D \setminus C(i)}(x^{(i)}) \right)^2 \quad \text{où } D \setminus C(i)$$

correspond à
l'ensemble des
données sans
compartiment
contenant

Etape #1 → sélectionner une

seule ensemble de caractéristiques $h_{\beta}^{D \setminus C(i)}$ → le modèle le dernier

→ pour chaque compartiment C_i entraîné sur $D \setminus C(i)$

$$h_{\beta}^{D \setminus C_i} = \arg \min_{\beta} \frac{1}{(M-1)N} \sum_{i \in D \setminus C_i} \left(t^{(i)} - h_{\beta}(x^{(i)}) \right)^2$$

