

Introduction à l'apprentissage

$$\mathcal{D} = \{(x^{(i)}, t^{(i)})\}_{i=1}^N$$

→ Regression linéaire

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

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

$$\tilde{x} = [1, \vec{x}]$$

$$\rightarrow \mathcal{L}(\beta) = \frac{1}{N} \sum_{i=1}^N (t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2$$

→ 2 approches → Descente de gradient
↳ Résolution des équations normales

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

$$\underline{X} \in \mathbb{R}^{N \times D+1}$$

→ Cas des caractéristiques dépendantes

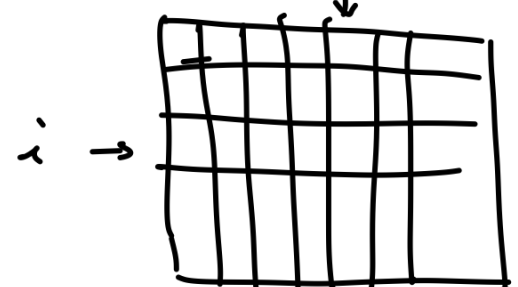
$$\rightarrow y(x) = \tilde{\underline{X}} \beta = \tilde{\underline{X}} (\tilde{\underline{X}}^T \tilde{\underline{X}})^{-1} \tilde{\underline{X}}^T t$$

Si caractéristiques linéairement dépendantes $\rightarrow \exists \vec{v} \quad \tilde{\underline{X}}^T \tilde{\underline{X}} v = 0$
 $\rightarrow \exists \lambda_i \approx 0$

pour rappel $\tilde{\underline{X}}^T \tilde{\underline{X}} = U \Lambda U^T$

$$(\tilde{\underline{X}}^T \tilde{\underline{X}})^{-1} = U \Lambda^{-1} U^T$$

$$\Lambda^{-1} = \begin{bmatrix} \lambda_1^{-2} & & & \\ & \ddots & & \\ & & \lambda_{D+1}^{-2} & \\ & & & \ddots \end{bmatrix}$$

$$\Lambda = \begin{bmatrix} \lambda_1 & 0 & \dots & \\ 0 & \lambda_2 & \dots & \\ \vdots & \vdots & \ddots & \vdots \\ 0 & & & \lambda_{D+1} \end{bmatrix}$$


f. caractéristiques sont dépendantes $\Rightarrow \exists \lambda_i \ll$

$$\Rightarrow (\underline{\tilde{X}}^T \underline{\tilde{X}})^{-1} \gg \underline{\beta}$$

$$\Rightarrow y(x) = \underline{\tilde{X}}_{k,t} (\underline{\tilde{X}}^T \underline{\tilde{X}})^{-1} \underline{\tilde{X}}^T \underline{t}$$
$$\underline{\tilde{X}}_{k,t} (\underline{\tilde{X}}^T \underline{\tilde{X}})^{-1} \underline{\tilde{X}}^T (t + \delta)$$

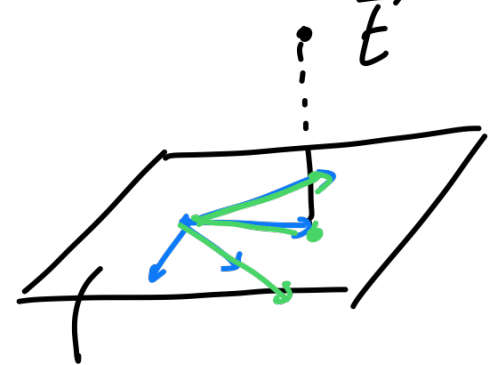
$\rightarrow$ On peut montrer que l'élément (i,j) de la matrice inverse $(\underline{\tilde{X}}^T \underline{\tilde{X}})^{-1}$ est donné par

$$\frac{(\underline{\varepsilon}_i)^T (\underline{\varepsilon}_j)}{\|\underline{\varepsilon}_i\|^2 \|\underline{\varepsilon}_j\|^2}$$

où $\underline{\varepsilon}_i = \underline{X}_i - \hat{\underline{X}}_i$ et $\hat{\underline{X}}_i$ est la projection de $\underline{x}_i$ (i.e. col. de $\underline{\tilde{X}}$) sur l'espace engendré par toutes les colonnes restantes.

On peut aussi voir la minimisation de la somme des carrés des résidus comme le calcul d'une projection orthogonale sur l'espace engendré par les colonnes de la matrice $\tilde{X}$

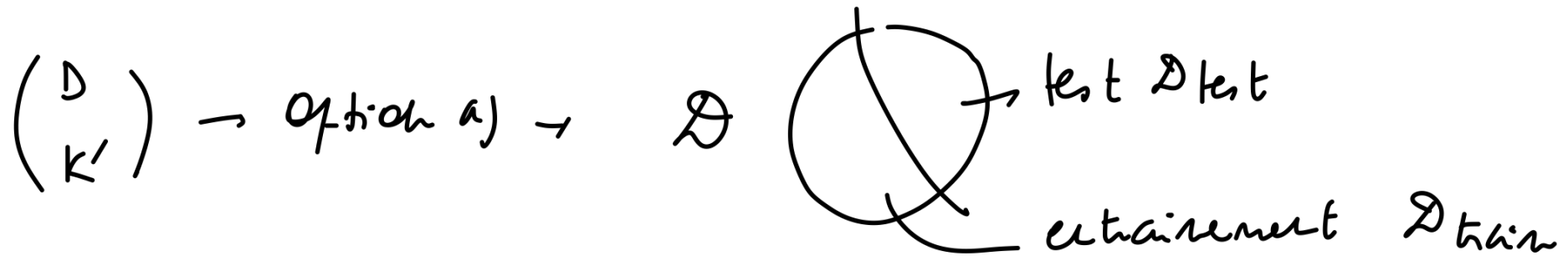
$$\min_{\beta} \| \vec{t} - \tilde{X} \beta \|^2$$



Espace engendré par les colonnes de $\tilde{X}$

Solution #1 $\rightarrow$ Gram Schmidt

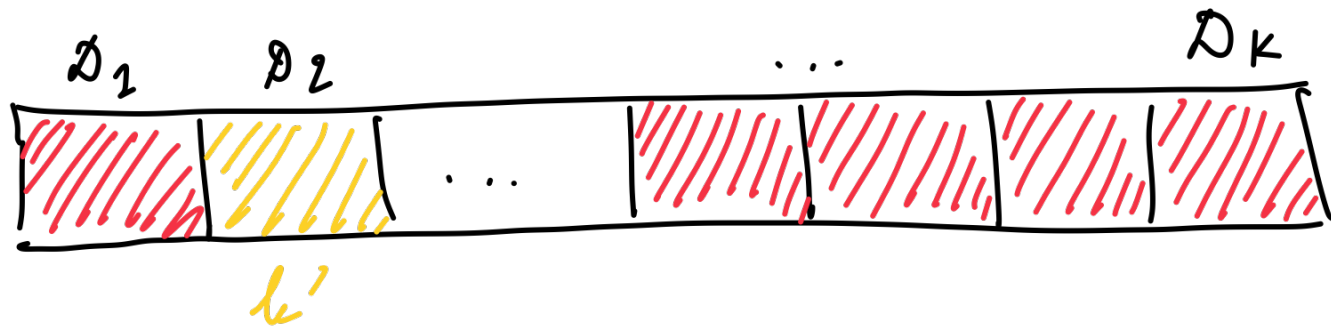
#2 $\rightarrow$ sélection du meilleur sous-ensemble



→ pour chaque sous ensemble parmi les $\binom{D}{K'}$ entraîner le modèle sur D_{train} et l'évaluer sur D_{test} .

Option b) Validation croisée à K compartments

pour chaque sous ensemble de caractéristique $\binom{D}{K'}$
on subdivise l'ensemble D en K compartments



$$D = \bigcup_{k=1}^K D_k$$

pour chaque compartiment k' , on entraîne le modèle sur les $K-1$ compartments restants

lorsque les modèles associés à chaque compartiment ont été entraînés, calculer l'erreur pour le modèle défini sur le sous-ensemble de caractéristiques

$$Err_{CV} = \frac{1}{N} \sum_{i=1}^N (t^{(i)} - h_{\beta}^{(i)}(x^{(i)}))^2$$

où $h_{\beta}^{(i)}(x)$ est le modèle entraîné sur tous les compartiments sauf celui contenant la donnée i

→ Caractéristiques des données → $h_{\beta}(x) = \beta_0 + \underbrace{\beta_1}_{\beta_1} x_1 + \underbrace{\beta_2}_{-\beta_1 + \Delta} x_2$

→ Alternative #3 **Penalti Ridge / Regression de type Ridge**

$$\min_{\beta} \quad \underset{\text{RIDGE}}{\ell(\beta)} = \frac{1}{N} \sum_{i=1}^N (t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2 + \lambda \sum_{j=1}^D \beta_j^2$$

Contrôle l'équilibre entre
fidélité aux données et complexité du modèle

→ Pour résoudre les équations normales, on commence par centrer les données

$$t^{(i)} \leftarrow t^{(i)} - \frac{1}{N} \sum_{i=1}^N t^{(i)} \quad \forall i$$

$$x_j^{(i)} \leftarrow x_j^{(i)} - \frac{1}{N} \sum_{i=1}^N x_j^{(i)} \quad \forall i, j$$

Si on calcule la dérivée par rapport à β_0 on trouve

$$\frac{\partial \mathcal{L}}{\partial \beta_0} = \frac{2}{N} \sum_{i=1}^N (t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)})) (-1)$$

Si on isole β_0 , on trouve

Si données centrées $\beta_0 = 0$

$$2\beta_0 = \frac{2}{N} \sum_{i=1}^N t^{(i)} - \frac{2}{N} \sum_{i=1}^N (\beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)})$$

On peut donc réécrire $\mathcal{L}_{\text{ridge}}(\beta)$ en fonction des $\beta_1, \dots, \beta_D$

uniquement

Si on suppose les données centrées

$$\mathcal{L}_{\text{ridge}}(\beta) = \frac{1}{N} \sum_{i=1}^N (t^{(i)} - (\beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2 + \lambda \sum_{j=1}^D \beta_j^2$$

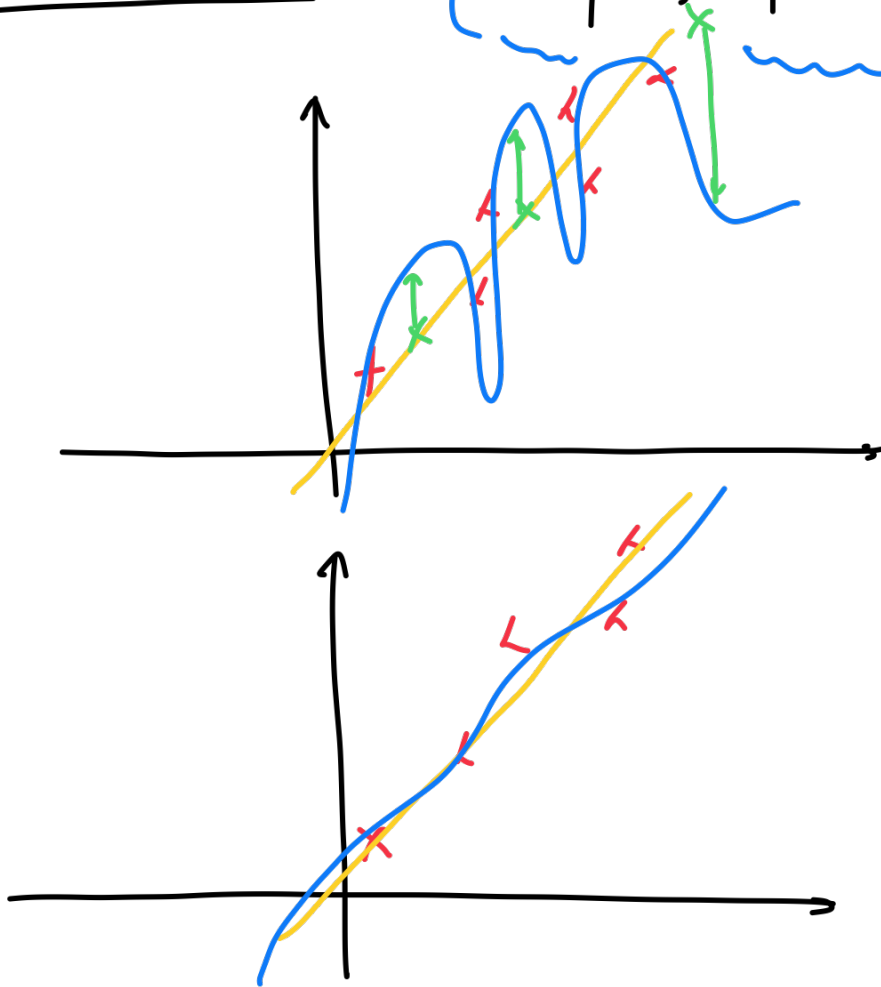
Illustration

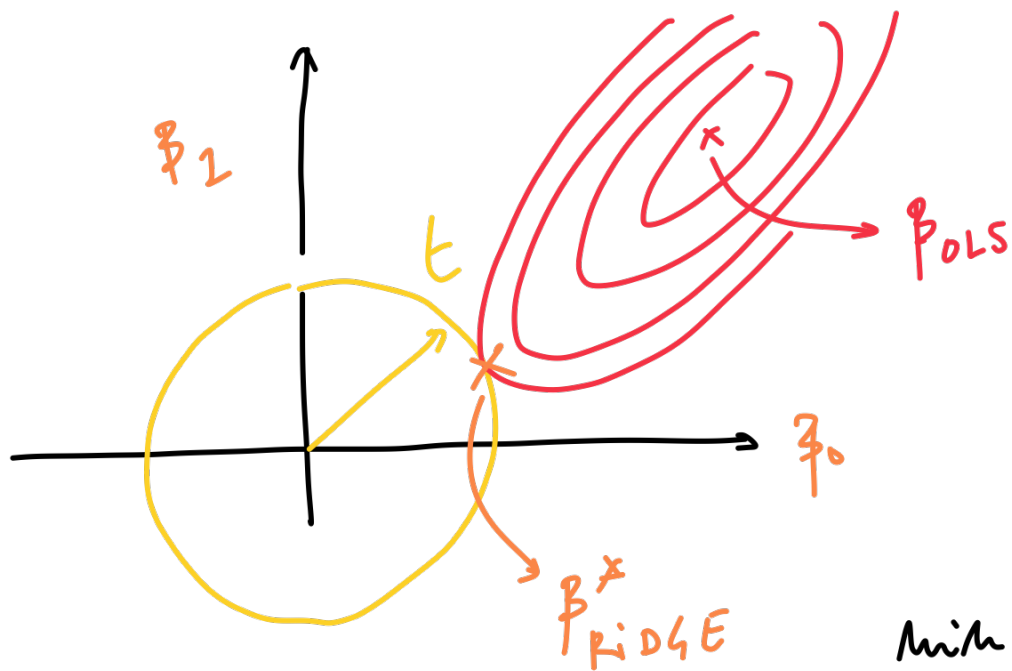
$$h_{\beta}(x) = \beta_0 + \beta_1 x + \beta_2 x^2 + \dots + \beta_D x^D$$

penalti de type RIDGE

$$h_{\beta}(x) = \beta_0 + \beta_1 x + \beta_2 x^2 + \dots + \beta_D x^D$$

$$\beta_1 \gg \beta_j \quad j \neq 1$$





Formulation ^{with} constraint

$$\min_{\beta} \frac{1}{N} \sum_{i=1}^N (t^{(i)} - h_{\beta}(x^{(i)}))^2 + \lambda \sum_{j=1}^D \beta_j^2$$

Formulation ^{with} constraint

$$\min_{\beta} \frac{1}{N} \sum_{i=1}^N (t^{(i)} - h_{\beta}(x^{(i)}))^2$$

s.t. $\sum_{j=1}^D \beta_j^2 \leq t$

