

Aujourd'hui: Motivations statistiques derrière Ridge et Lasso

→ Equilibre biais variance

→ Classification

→ Regression linéaire

$$y(x) = \beta_0 + \beta_1 x_1 + \dots + \beta_D x_D \quad \vec{x} \in \mathbb{R}^D, \{ (x_i, t^{(i)}) \}_{i=1}^N$$

Pour rappel, dans le cadre de la minimisation RSS

On avait fait les hypothèses suivantes

$$t^{(i)} \sim \beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)} + \varepsilon^{(i)} \quad \varepsilon^{(i)} \sim \mathcal{N}(0, \sigma^2)$$

$\varepsilon^{(i)}$ indépendants et identiquement distribués (i.i.d)

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

$$p(\{t^{(i)}, \{x_{i=1}^D\} | \beta) = \prod_{i=1}^N \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2}{2\sigma^2}\right)$$

→ Estimateur de Maximum de vraisemblance

$$\beta_{MLE} = \underset{\beta}{\operatorname{argmax}} p(\{t^{(i)}, \{x_{i=1}^D\} | \beta)$$

En particulier, on a montré $\beta_{MLE} = \beta_{RSS}$

Théorème de Bayes

$$P(B | A) = \frac{P(B \cap A)}{P(A)} = \frac{P(A | B) P(B)}{P(A)}$$

Estimateur de Maximum a Posteriori β_{MAP}

au lieu d'estimer le modèle en maximisant $P(\{t^{(i)}\}_{i=1}^N | \beta)$

on peut chercher le modèle qui maximise $P(\beta | \{t^{(i)}, x^{(i)}\}_{i=1}^N)$

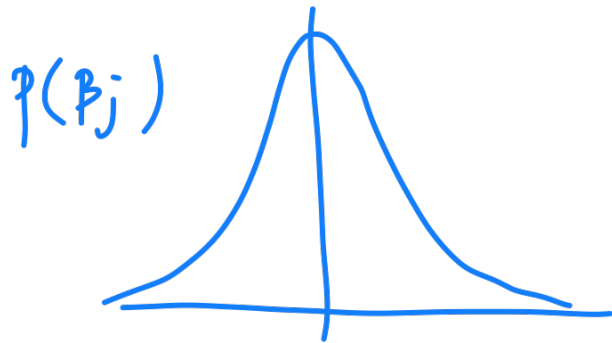
En utilisant Bayes, on trouve

$$P(\underline{\beta} | \{t^{(i)}, x^{(i)}\}_{i=1}^N) = \frac{P(\{t^{(i)}\}_{i=1}^N | \beta, \{x^{(i)}\}_{i=1}^N) P(\beta)}{P(\{t^{(i)}\}_{i=1}^N)}$$

$$p(\{t^{(i)}\}_{i=1}^N) = \sum_{\beta_k} p(\{t^{(i)}\}_{i=1}^N | \beta_k) p(\beta_k) \rightarrow \text{independent de } \beta$$

$$\beta_{MAP} = \operatorname{argmax}_{\beta} p(\{t^{(i)}\}_{i=1}^N | \beta, \{x^{(i)}\}_{i=1}^N) p(\beta)$$

Exemples de choix pour $p(\beta)$: $\rightarrow p(\beta) \sim \text{Uniforme}$



$$\beta_{MAP} = \beta_{MLE}$$

$\rightarrow p(\beta) \sim \text{Gaussienne}$

$$p(\beta) \sim \prod_{j=1}^D \frac{1}{\sqrt{2\pi}\xi} \exp\left(-\frac{\beta_j^2}{2\xi^2}\right)$$

β_j indépendamment et indépend.
distribués selon une
loi Normale

→ Pour un a priori Gaussien sur les β_j l'estimateur β_{MAP}

se réécrit

$$\beta_{MAP} = \arg \max_{\beta} \prod_{i=1}^W \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2}{2\sigma^2}\right) \prod_{j=1}^D \frac{1}{\sqrt{2\pi}\xi} \exp\left(-\frac{\beta_j^2}{2\xi^2}\right)$$

On commence par appliquer le logarithme ce qui donne

$$\beta_{MAP} = \arg \max_{\beta} \sum_{i=1}^W \log\left(\frac{1}{\sqrt{2\pi}\sigma}\right) - \frac{(t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2}{2\sigma^2}$$

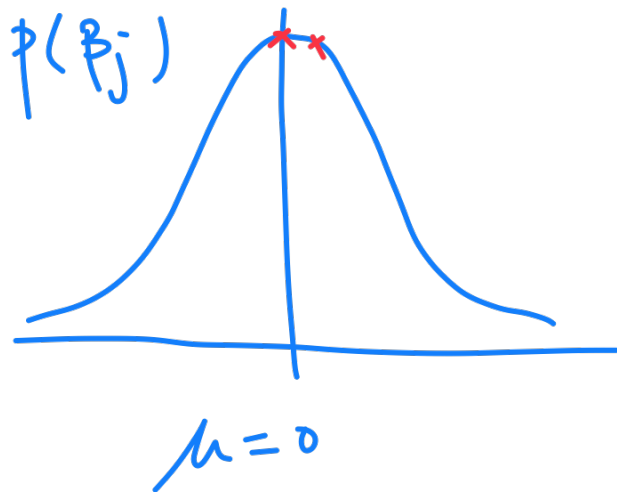
$+ \sum_{j=1}^D \log\left(\frac{1}{\sqrt{2\pi}\xi}\right) - \frac{\beta_j^2}{2\xi^2}$

independants de β

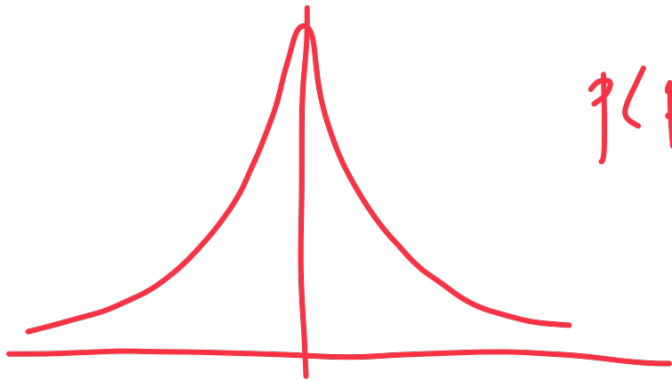
$$\beta_{\text{HAP}} = \arg \max_{\beta} - \sum_{i=1}^N \frac{(t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2}{2\sigma^2} - \sum_{j=1}^D \frac{\beta_j^2}{2\xi^2}$$

$$= \arg \min_{\beta} \sum_{i=1}^N \frac{(t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2}{2\sigma^2} + \sum_{j=1}^D \frac{\beta_j^2}{2\xi^2}$$

$\beta_{\text{RIDGE}} = \beta_{\text{HAP}}$ avec un a priori gaussien sur les β_j



→ Pour favoriser davantage l'annulation des coefficients on peut choisir un a priori de type Laplace



$$p(\beta_j) = \frac{1}{2\lambda} \exp\left(-\frac{|\beta_j|}{\lambda}\right)$$

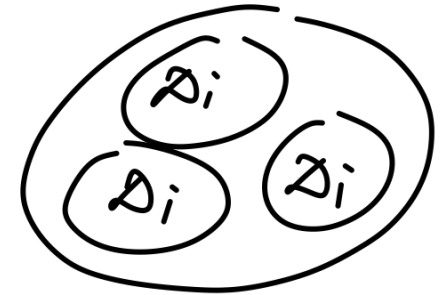
$$\beta_{MAP} = \arg \max_{\beta} \prod_{i=1}^N \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(\ell^{(i)} - h_{\beta}(x^{(i)}))^2}{2\sigma^2}\right) \prod_{j=1}^D \frac{1}{2\lambda} \exp\left(-\frac{|\beta_j|}{\lambda}\right)$$

en appliquant le log et en transformant la maximisation en minimisation, on retrouve

$$\beta_{\text{HAP}} = \underset{\beta}{\text{argmin}} \underbrace{\sum_{i=1}^N \frac{(t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}))^2}{2\sigma^2}}_{\text{LSE}} + \sum_{j=1}^D \frac{|\beta_j|}{\xi}$$

Équilibre biais - variance

$(x, t(x))$
donnée test



Erreur quadratique moyenne

$$\text{MSE} = \mathbb{E}_{D_i} \{ (h_{D_i}(x) - t(x))^2 \}$$

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

$$= \mathbb{E}_{D_i} \{ (h_{D_i}(x) - \mathbb{E}_{D_i} h_{D_i}(x) + \mathbb{E}_{D_i} h_{D_i}(x) - t(x))^2 \}$$

En appliquant le lemme, on obtient

$$\begin{aligned}
 \text{MSE}(x) &= \mathbb{E}_{\mathcal{D}_i} \{ (\underbrace{h_{\mathcal{D}_i}(x)}_{\text{red wavy}} - \underbrace{\mathbb{E}_{\mathcal{D}_i} h_{\mathcal{D}_i}(x)}_{\text{red wavy}})^2 \} + \cancel{\mathbb{E}_{\mathcal{D}_i} \{ (\underbrace{\mathbb{E}_{\mathcal{D}_i} h_{\mathcal{D}_i}(x)}_{\text{blue wavy}} - t(x))^2 \}}^{\text{deterministic}} \\
 &\quad + \underbrace{2 \mathbb{E}_{\mathcal{D}_i} \{ (h_{\mathcal{D}_i}(x) - \mathbb{E}_{\mathcal{D}_i} h_{\mathcal{D}_i}(x)) (\mathbb{E}_{\mathcal{D}_i} h_{\mathcal{D}_i}(x) - t(x)) \}}_{=0}
 \end{aligned}$$

$$\text{MSE}(x) = \underbrace{\mathbb{E}_{\mathcal{D}} \{ (h_{\mathcal{D}_i}(x) - \mathbb{E}_{\mathcal{D}_i} h_{\mathcal{D}_i}(x))^2 \}}_{\text{Variance}} + \underbrace{(\mathbb{E}_{\mathcal{D}_i} h_{\mathcal{D}_i}(x) - t(x))^2}_{\text{bias}^2}$$

