

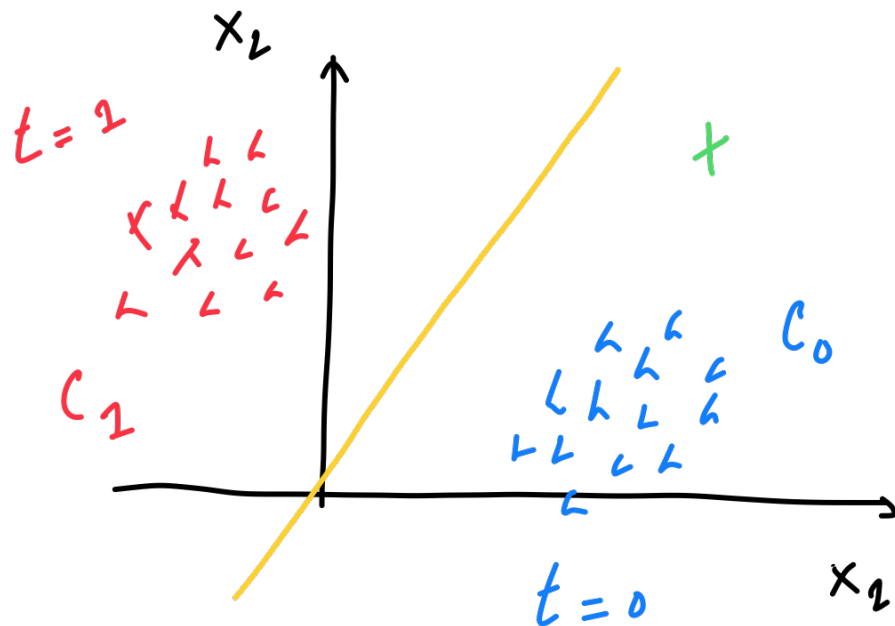
Apprentissage supervisé

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

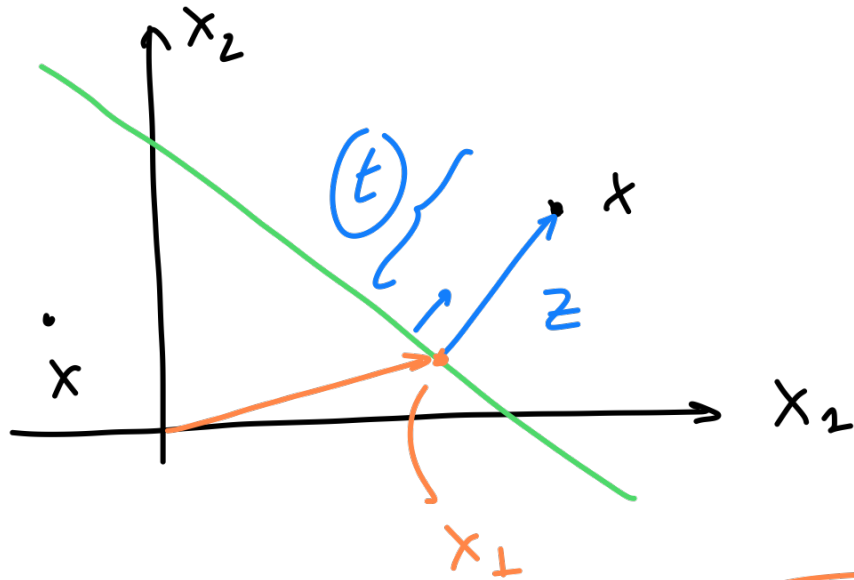
$$\rightarrow h_{\beta}(x) = \beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)}$$

Partie 2 : Classification

$$t^{(i)} \in \{1, 2, \dots, K\}$$



$$t^{(i)} \in \{0, 1\}$$



$$\vec{x} = \vec{x}_{\perp} + \vec{z}$$

$$\downarrow$$

$$(x_1, x_2) = (x_{\perp})_1, (x_{\perp})_2 + (z_1, z_2)$$

→ $x_{\perp}$ est sur le plan : $\beta_0 + \beta_1(x_{\perp})_1 + \beta_2(x_{\perp})_2 = 0$ ←

$$\beta_0 + \beta_1 x_1 + \beta_2 x_2 = \underbrace{\beta_0 + \beta_1(x_{\perp})_1 + \beta_2(x_{\perp})_2}_{=0} + \underbrace{\beta_1 z_1 + \beta_2 z_2}_{t \|(\beta_1, \beta_2)\|} \geq 0$$

$$\vec{z} = \frac{t(\beta_1, \beta_2)}{\|(\beta_1, \beta_2)\|} \leftarrow$$

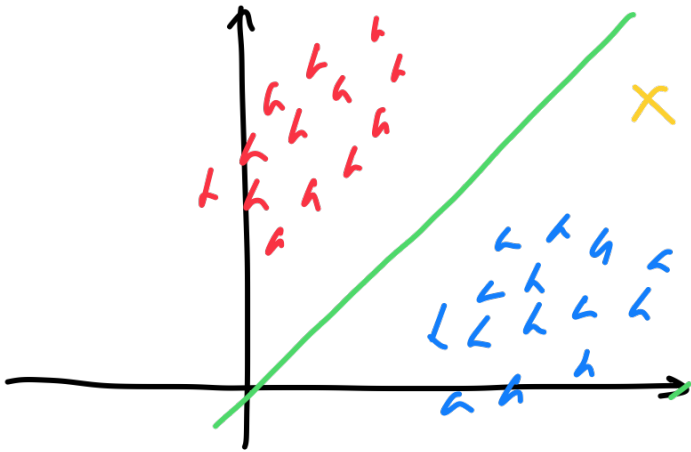
$$\langle \vec{z}, (\beta_1, \beta_2) \rangle = \frac{t \|(\beta_1, \beta_2)\|^2}{\|(\beta_1, \beta_2)\|}$$

Function cost

min
 β

$$l(\beta) = \frac{1}{N} \sum_{i=1}^N (t^{(i)} - (\beta_0 + \beta_1 x_1^{(i)} + \beta_2 x_2^{(i)}))^2$$

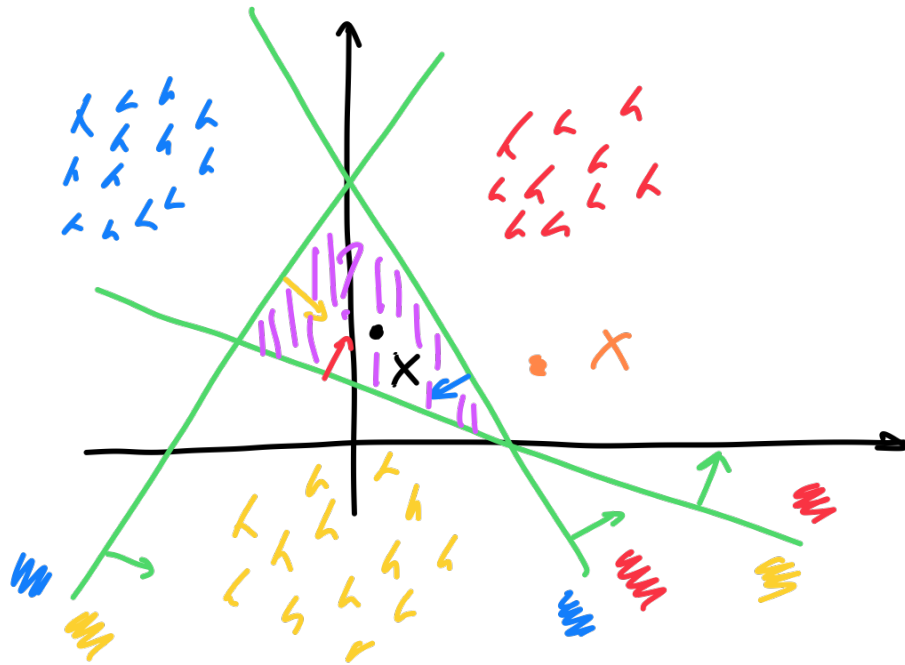
$\in \{1, 0\}$



$$\text{prediction}(x) = \beta_0 + \beta_1 x_1 + \beta_2 x_2$$

C_1 C_0

if $\text{prediction}(x) > \frac{1}{2} \rightarrow C_1$
 $< \frac{1}{2} \rightarrow C_0$

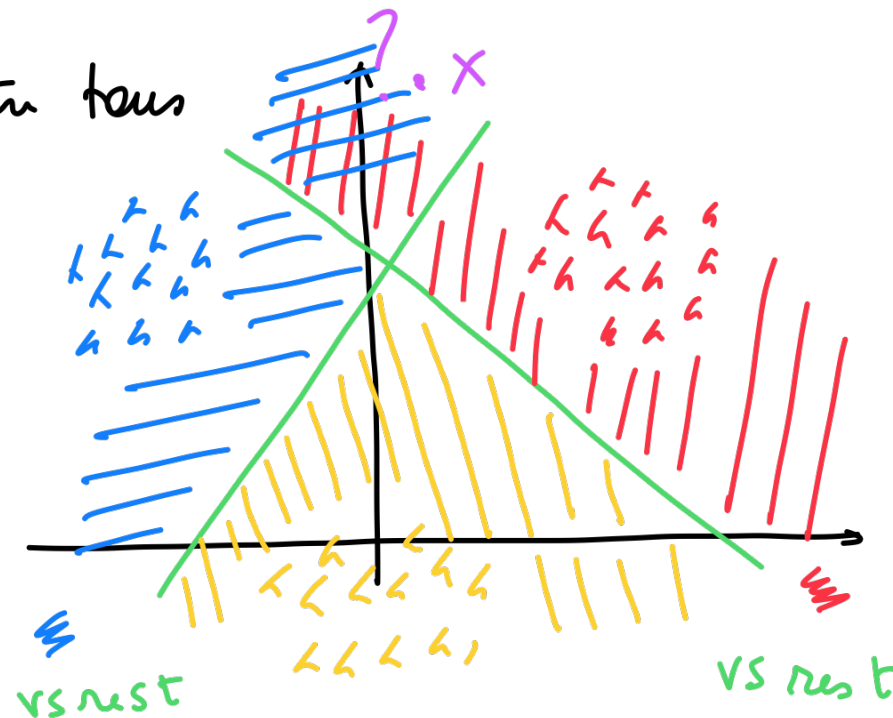


→ Approche 1 contre 1

$$\frac{K(K-1)}{2} \text{ discriminants}$$

Classe de X donnée par vote à la majorité

Approche 1 contre tous



$(K-1)$
discriminants

→ Solution #3 : Discriminant à K classes

$$t^{(i)} \in \{0, 1\} \rightarrow \vec{t}^{(i)} = [0, 0, \dots, \underset{\substack{\uparrow \\ K}}{1}, 0, 0, \dots, 0] \text{ (one hot encoding)}$$

$$\vec{B} = \begin{bmatrix} \beta_0^{(2)} & \dots & \beta_D^{(2)} \\ \beta_0^{(2)} & \dots & \beta_D^{(2)} \\ \vdots & & \\ \beta_0^{(K)} & \dots & \beta_D^{(K)} \end{bmatrix} \in \mathbb{R}^{K \times (D+1)}$$

discriminant pour la classe K

$$\text{Si } x \in \text{classe } K \Rightarrow \beta_0^{(K)} + \beta_1^{(K)} x_1 + \dots + \beta_D^{(K)} x_D \approx 1$$

$$\text{Si } x \notin \text{classe } K \Rightarrow \beta_0^{(K)} + \beta_1^{(K)} x_1 + \dots + \beta_D^{(K)} x_D \approx 0$$

$$\tilde{X} = \begin{bmatrix} 1 & x_1^{(2)} & \dots & x_D^{(2)} \\ \vdots & \vdots & & \vdots \\ 1 & x_1^{(N)} & & x_D^{(N)} \end{bmatrix} \in \mathbb{R}^{N \times (D+1)}$$

$$\begin{bmatrix} \beta_0^{(2)} & \dots & \beta_0^{(K)} \\ \vdots & & \vdots \\ \beta_D^{(2)} & & \beta_D^{(K)} \end{bmatrix}$$

$$\tilde{X} B^T \approx \underbrace{\begin{bmatrix} 1 & 0 & \dots & 0 \\ 0 & \dots & 1 & \dots \end{bmatrix}}_{N \times K}$$

$$T = \begin{bmatrix} -\vec{t}^{(2)}- \\ -\vec{t}^{(2)}- \\ \vdots- \\ -\vec{t}^{(N)}- \end{bmatrix}$$

argmin $\underset{B}{B}$ $\frac{1}{NK} \sum_{i=1}^N \sum_{k=1}^K (t_k^{(i)} - (\beta_0^{(k)} + \beta_1^{(k)} x_1^{(i)} + \dots + \beta_D^{(k)} x_D^{(i)}))^2$

argmin $\underset{B}{B}$ $\frac{1}{NK} \langle \underbrace{T - \tilde{X} B^T}_{=}, \underbrace{T - \tilde{X} B^T}_{=} \rangle_F$

↓

$$\langle A, B \rangle_F = \sum_{i,j} A_{ij} B_{ij} \quad \langle B^T, \tilde{X}^T (T - \tilde{X} B^T) \rangle$$

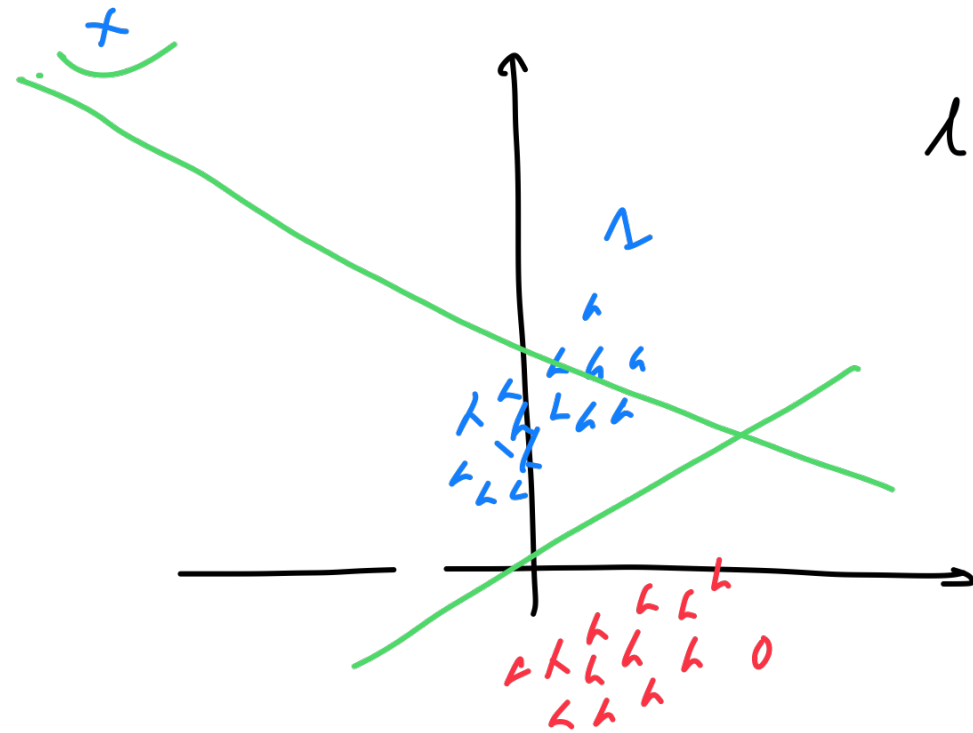
$$\underline{\underline{\bar{B} = (\tilde{X}^T \tilde{X})^{-1} \tilde{X}^T T}}$$

Soit x une donnée test

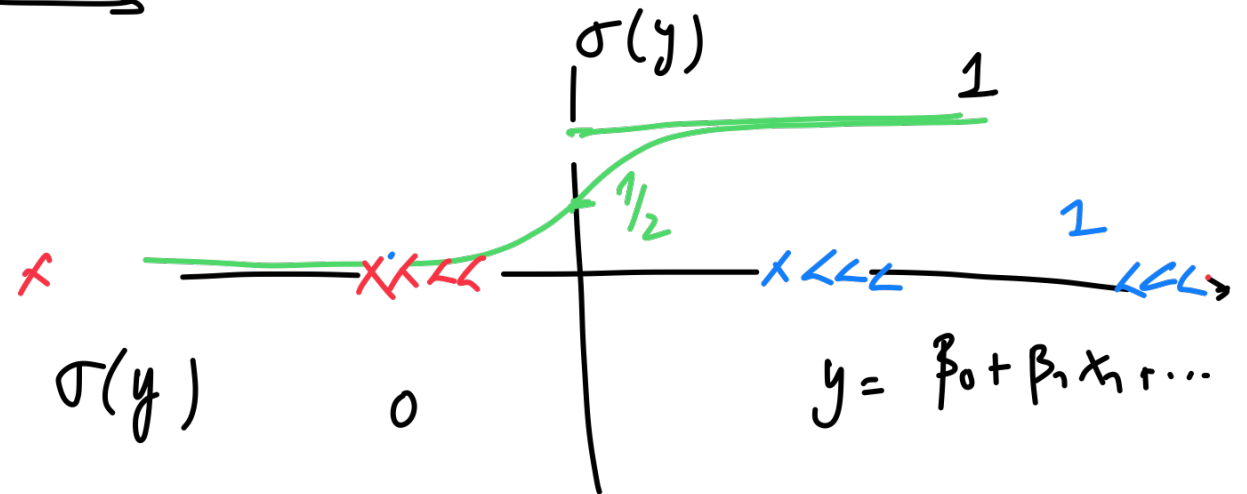
→ la classe de x est calculée via

$$\begin{bmatrix} \beta_0^{(2)} \dots \beta_D^{(2)} \\ \vdots \\ \beta_0^{(K)} \dots \beta_D^{(K)} \end{bmatrix} \begin{bmatrix} 1 \\ x_2 \\ \vdots \\ x_D \end{bmatrix} = \begin{bmatrix} t_2(x) \\ \vdots \\ t_K(x) \end{bmatrix}$$

$$\text{Classe de } x = \underset{k}{\operatorname{argmax}} \quad t_k(x)$$



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



→ Fonction d'activation $\sigma(y)$

$$y(x) = \sigma(\beta_0 + \beta_1 x_1 + \dots + \beta_D x_D)$$

$$\sigma(y) = \frac{1}{1 + e^{-y}}$$

$$p(t = \text{Yes} \mid x) = \sigma(\beta_0 + \beta_1 x_1 + \dots + \beta_D x_D)$$

$$p(t = \text{No} \mid x) = 1 - \sigma(\beta_0 + \beta_1 x_1 + \dots + \beta_D x_D)$$