

2023-2024:

$$P(X=x) = \lambda e^{-\lambda x}$$

$$E(X) = 1/\lambda$$

$N = 21067$ images ($L \times L = 50$ pixel), $K=3$ classes

1) dimension D des données d'entrée: $L^2 = 50^2 = 2500$

Taille de $X: (N, D)$ ($X \in \mathbb{R}_+^{N \times D}$)

Taille de $Y: N$

$$(Y \in \{1, 2, 3\}^N)$$

ou $\{ \text{rugby, volley, foot} \}^N$

2) On s'intéresse au $\frac{1}{L^2}$ pixel des images de classe $K=1$

a) modèle associé? param à prendre? $\int \prod_{i=1}^N y_n = 1$. Chaque $x_{n,1}$ est une valeur de gris (contour, ≥ 0)

$\Rightarrow$ On va modéliser le pixel 1 ($x_{n,1}$) de la classe 1 comme issu d'une loi exponentielle

de paramètre $\lambda_{n,1} = P(X_1=x_1) = \lambda e^{-\lambda x_1}$ (abus de notation)

$$X_1 | (Y=1) \sim \text{Exp}(\lambda_{1,1})$$

b) estimation du maximum de vraisemblance:

$\lambda_{\#}$ contrôle l'échelle: λ grand $\Rightarrow$ valeurs typiques petites

λ petit $\Rightarrow$ valeurs typiques grandes

et on voit que pour une exponentielle, $E[X] = 1/\lambda$

$\Rightarrow$ on peut dériver $1/\lambda_{n,1}$ sera la valeur moyenne du pixel 1 pour les images de classe $K=1$

3) paramètres du modèle global? (θ ?)

On fait du Naïf Bayes exponentiel:

- on a un modèle par classe $K \in \{1, \dots, K\}$

- dans chaque classe, chaque pixel $j \in \{1, \dots, L^2\}$ est modélisé par une exponentielle avec son propre λ

(V classe 1 et Y_j pixel, $X_j | (Y=K) \sim \text{Exp}(\lambda_{K,j})$)

Donc $\theta = \{ \lambda_{K,j}, K=1, \dots, K, j=1, \dots, L^2 \}$ (il y en a KL^2 : ici $3 \times 2500 = 7500$ paramètres)

$\Rightarrow$ Si on considère la fréquence d'apparition de chaque classe comme un paramètre à apprendre: on rajoute $K=3$ paramètres de plus.

4) Maximiser la vraisemblance des données:

[Notations: $X = (X_1, \dots, X_N$ la v.a. associée à tout le dataset, $(\tilde{x}_n)_n$ l'ensemble des données effectivement vues, $(y_n)_n$ les labels effectivement vus]

$\Rightarrow$ Cela revient à apprendre les valeurs des paramètres θ telles que $\mathcal{L}(\theta) = P(X = (\tilde{x}_n) | Y = (y_n), \theta)$ soit maximale.

On cherche $\theta^* = \arg \max_{\theta} (\mathcal{L}(\theta))$

$$\theta^* = \arg \max_{\theta} (\log(\mathcal{L}(\theta))) \quad [\log \uparrow]$$

$$= \arg \max_{\theta} (\log(P(X = (\tilde{x}_n) | Y = (y_n), \theta))) \quad [\text{def vraisemblance: } \mathcal{L}(\theta) = P(\text{données} | \theta)]$$

$$= \arg \max_{\theta} (\log(\prod_{n=1}^N P(X_n = \tilde{x}_n | Y_n = y_n, \theta))) \quad [\text{pour indépendance des } X_n, n \in [1, N]]$$

$$= \arg \max_{\theta} \left(\sum_{n=1}^N \log(P(X_n = \tilde{x}_n | Y_n = y_n, \theta)) \right) \quad [\log(\prod_{n=1}^N a_n) = \sum_{n=1}^N \log a_n]$$

or $Y_n, Y_K \quad P(X_n = \tilde{x}_n | Y_n = K, \theta) = \prod_{d=1}^D P(X_{n,d} = x_{n,d} | Y_n = K, \theta)$ [hypothèse Naïf Bayes: indépendance conditionnelle des pixels sachant $d=1 \quad Y_n = K$. Abus de notation: $X_{n,d}$ et la v.a. et $x_{n,d}$ la valeur prise dans la réalisation de cette v.a. que l'on nomme $x_{n,d}$]

$$= \prod_{d=1}^D \lambda_{K,d} e^{-\lambda_{K,d} x_{n,d}} \prod_{Y_n=K}$$

le modèle de la distribution des features d'entrée est Exponentiel. (les θ sont les $\lambda_{K,d}$)

$$\begin{aligned} \text{donc } \theta^* &= \arg \max_{\theta} \left(\sum_{n=1}^N \log \left(\prod_{d=1}^D \lambda_{K,d} e^{-\lambda_{K,d} x_{n,d}} \right) \prod_{Y_n=K} \right) \\ &= \arg \max_{\theta} \left(\sum_{n=1}^N \sum_{d=1}^D \log(\lambda_{K,d} e^{-\lambda_{K,d} x_{n,d}}) \prod_{Y_n=K} \right) \\ &= \arg \max_{\theta} \left(\sum_{n=1}^N \sum_{d=1}^D [\log(\lambda_{K,d}) - \lambda_{K,d} x_{n,d}] \prod_{Y_n=K} \right) \end{aligned}$$

$\Rightarrow$ On trouve le maximum de cette fonction [on dérive par rapport à θ , i.e. $\lambda_{K,d}$ et on cherche la valeur de $\lambda_{K,d}$ telle que la dérivée soit nulle]

$$0 = \frac{\partial \log \mathcal{L}(\theta)}{\partial \lambda_{K,d}} = \frac{\partial}{\partial \lambda_{K,d}} \left(\sum_{n=1}^N \sum_{d=1}^D [\log(\lambda_{K,d}) - \lambda_{K,d} x_{n,d}] \prod_{Y_n=K} \right)$$

$\hookrightarrow$ pour $d \neq d'$
 $\log \lambda_{K,d} - \lambda_{K,d} x_{n,d}$ est de $\% \lambda_{K,d}$
et le terme dans la somme vaut 0 si $d \neq d'$

$$\Rightarrow 0 = \sum_{n=1}^N \left[\frac{1}{\lambda_{K,d}} - x_{n,d} \right] \prod_{Y_n=K} \prod_{d' \neq d}$$

$$\Rightarrow \sum_{n=1}^N \frac{1}{\lambda_{K,d}} \prod_{Y_n=K} = \sum_{n=1}^N x_{n,d} \prod_{Y_n=K}$$

$$\Rightarrow \frac{1}{\lambda_{K,d}} \sum_{n=1}^N \prod_{Y_n=K} = \sum_{n=1}^N x_{n,d} \prod_{Y_n=K}$$

$$\Rightarrow \frac{1}{\lambda_{K,d}} = \frac{\sum_{n=1}^N x_{n,d} \prod_{Y_n=K}}{\sum_{n=1}^N \prod_{Y_n=K}} \quad \Rightarrow \text{moyenne des pixels numéro } d \text{ des images de classe } K$$

$\hookrightarrow$ compte le nombre d'images dans la classe

5) Comment visualiser ces données?

On a appris des param $\lambda_{K,d}$ (un par classe K et par pixel d). Comme chaque $\lambda_{K,d}$ correspond directement à un pixel, on peut les regrouper en forme $L \times L$. Donc pour chaque classe K , je construis une "image des paramètres" $\lambda_{K,d}$ ($\lambda_{K,d}$ est $\lambda_{K,d}(c,d)$) (mod $L \times L$), je l'affiche en niveaux de gris (ou en heatmap)

Interprétation: $E[X_{K,d}] = 1/\lambda_{K,d}$: afficher $1/\lambda_{K,d}$ donne une image qui ressemble à la "moyenne" des images de classe K .

6) Ce modèle est aussi un modèle génératif. Comment échantillonner de nouvelles images à partir du modèle?

On choisit un K . $\hookrightarrow$ (définir une loi de probas sur les images conditionnellement à une classe K)

On génère $\tilde{x} \in \mathbb{R}_+^D$ en tirant, pour chaque pixel $d=1, \dots, D$, indépendamment: $x_d \sim \text{Exp}(\lambda_{K,d})$, et on remet $\tilde{x}$ en forme $L \times L$.

$\Rightarrow$ Ces images se ont facilement identifiables comme fausses car le modèle impose l'indépendance des pixels. Il n'y a donc pas de corrélations spatiales (pixels voisins). On obtient donc du "bruit" et des variations brutales $\tilde{m}$ dans des zones censées être uniformes.

7) [Bonus] fonction de décision

Pour un nouveau point de donnée, on connaît $\vec{x}$ et on souhaite prédire y , connaissant les paramètres θ .

On choisit la classe k comme celle qui correspond le plus probablement à l'image observée $\vec{x}$,

$$k^* = \arg \max_{1 \leq k \leq K} \left(\mathbb{P}(y=k | \vec{x}, \theta) \right)$$

$$= \arg \max_{1 \leq k \leq K} \left(\frac{\mathbb{P}(y=k | \vec{x}, \theta)}{\mathbb{P}(\vec{x} | \theta)} \right) \quad \text{Bayes: } \mathbb{P}(A|B, C) \mathbb{P}(B|C) = \mathbb{P}(A, B|C)$$

$$= \arg \max_{1 \leq k \leq K} \left(\frac{\mathbb{P}(\vec{x} | y=k, \theta) \mathbb{P}(y=k | \theta)}{\mathbb{P}(\vec{x} | \theta)} \right)$$

$$= \arg \max_{1 \leq k \leq K} \left(\mathbb{P}(\vec{x} | y=k, \theta) \right) \mathbb{P}(y=k | \theta)$$

Les probabilités $\mathbb{P}(\vec{x} | y=k, \theta)$ et $\mathbb{P}(y=k | \theta)$ sont des probabilités d'observer une image $\vec{x}$ appartenant à la classe k , connaissant θ mais pas $\vec{x}$.

$$\forall k, \pi_k = \mathbb{P}(y=k | \theta) = \frac{\sum_{n=1}^N \mathbb{1}_{y_n=k}}{N}$$

Proportion d'images de la classe k dans le training set

$$\text{donc } k^* = \arg \max_{1 \leq k \leq K} \left(\mathbb{P}(\vec{x} | y=k, \theta) \mathbb{P}(y=k | \theta) \right) = \arg \max_{1 \leq k \leq K} \left(\prod_{d=1}^D \lambda_{k,d} e^{-\lambda_{k,d} x_d} \frac{\sum_{n=1}^N \mathbb{1}_{y_n=k}}{N} \right)$$

$$= \arg \max_{1 \leq k \leq K} \left(\log \left(\prod_{d=1}^D \lambda_{k,d} e^{-\lambda_{k,d} x_d} \right) - \log \frac{\sum_{n=1}^N \mathbb{1}_{y_n=k}}{N} \right)$$

$$= \arg \max_{1 \leq k \leq K} \left(\sum_{d=1}^D [\log \lambda_{k,d} - \lambda_{k,d} x_d] + \log \pi_k \right)$$

2021: Bayésien Naïf Gaussien

$$\mathbb{P}(X_n \in [x, x+dx]) = p(x) dx$$

$$p(x) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$$

$$\text{abus de notation: } \mathbb{P}(X_n = x) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$$

1) $\theta = \{?, ?\}$

$$\text{ici, } x_d | (y_k = a) \sim \mathcal{N}(\mu_{k,d}, \sigma_{k,d}^2)$$

$$\text{donc } \theta = \{\mu_{k,d}, \sigma_{k,d}^2, k=1, \dots, K, d=1, \dots, D\} \cup \{\pi_k, k=1, \dots, K\}$$

2) $|\theta| = 2KD + 2$

3) Écrire l'algorithme d'apprentissage (estimation du maximum de vraisemblance) accorde des votes

$$\text{Mais: } \mathbb{P}(X_n = \vec{x} | y_n = k, \theta) = \prod_{d=1}^D \frac{1}{\sqrt{2\pi\sigma_{k,d}^2}} \exp\left(-\frac{(x_{n,d} - \mu_{k,d})^2}{2\sigma_{k,d}^2}\right)$$

$$\text{Les @ étant les } \mu_{k,d}, \sigma_{k,d}^2$$

$$\theta^* = \arg \max_{\theta} \left(\sum_{n=1}^N \log \left(\prod_{d=1}^D \frac{1}{\sqrt{2\pi\sigma_{k,d}^2}} \exp\left(-\frac{(x_{n,d} - \mu_{k,d})^2}{2\sigma_{k,d}^2}\right) \right) \mathbb{1}_{y_n=k} \right)$$

$$= \arg \max_{\theta} \left(\sum_{n=1}^N \sum_{d=1}^D \log \left(\frac{1}{\sqrt{2\pi\sigma_{k,d}^2}} \exp\left(-\frac{(x_{n,d} - \mu_{k,d})^2}{2\sigma_{k,d}^2}\right) \right) \mathbb{1}_{y_n=k} \right)$$

$$= \arg \max_{\theta} \sum_{n=1}^N \sum_{d=1}^D \left(-\frac{1}{2} \log(2\pi\sigma_{k,d}^2) + \left[-\frac{(x_{n,d} - \mu_{k,d})^2}{2\sigma_{k,d}^2} \right] \right) \mathbb{1}_{y_n=k}$$

$$\hookrightarrow 0 = \frac{\partial \log \mathcal{L}(\theta)}{\partial \mu_{k,d}} = 0 + \left(\sum_{n=1}^N \sum_{d=1}^D \left[-\frac{1}{\sigma_{k,d}^2} \left[-\frac{(x_{n,d} - \mu_{k,d})^2}{2\sigma_{k,d}^2} \right] \right) \mathbb{1}_{y_n=k} \right)$$

$$\Rightarrow 0 = \left(\sum_{n=1}^N \left[\frac{1}{\sigma_{k,d}^2} (x_{n,d} - \mu_{k,d}) \right] \mathbb{1}_{y_n=k} \right) + \sum_{d \neq d'}^D \left(\sum_{n=1}^N [x_{n,d} - \mu_{k,d}] \mathbb{1}_{y_n=k} \right)$$

$$\Rightarrow \sum_{n=1}^N \mu_{k,d} \mathbb{1}_{y_n=k} = \sum_{n=1}^N x_{n,d} \mathbb{1}_{y_n=k} \Rightarrow \mu_{k,d} = \frac{\sum_{n=1}^N x_{n,d} \mathbb{1}_{y_n=k}}{\sum_{n=1}^N \mathbb{1}_{y_n=k}} \quad \text{ce qui correspond à faire la moyenne de la feature } n^{\circ} d \text{ des points de classe } k$$

4) **Représentation visuelle des paramètres appris par le modèle**

$\mu_{k,d}$: pour chaque classe k , remettre $(\mu_{k,d})_{d=1, \dots, D}$ en matrice $L \times D$ et l'afficher en niveau de gris (image moyenne de la classe)

$\sigma_{k,d}$ (ou $\sigma_{k,d}^2$): pour chaque classe k , afficher $(\sigma_{k,d})_d$ en image $L \times D$ (carte d'écart-type/variance par pixel)

π_k : afficher le vecteur $(\pi_1, \dots, \pi_K)$ en diagramme de barres (une barre par classe k par $k \geq 2$ avec $\sum_{k=1}^K \pi_k = 1$)

5) $\vec{x}$? échantillonner de nouvelles images

Il suffit de générer des images en tirant au hasard la valeur des pixels (valeur du niveau de gris) selon la loi (gaussienne) de chaque pixel, pour chaque indépendamment des autres pixels

6) **Fonction de décision**

puis pour la Gaussienne:

$$k^* = \arg \max_k \prod_{d=1}^D \frac{1}{\sqrt{2\pi\sigma_{k,d}^2}} \exp\left(-\frac{(x_d - \mu_{k,d})^2}{2\sigma_{k,d}^2}\right) = \arg \max_k \left[\log \prod_{d=1}^D \frac{1}{\sqrt{2\pi\sigma_{k,d}^2}} - \frac{(x_d - \mu_{k,d})^2}{2\sigma_{k,d}^2} \right]$$

2024 - 2025: Gaussienne (+ qq's questions)

$$x_0 \in \mathbb{R}, X = (x_1, x_2, \dots, x_N)$$

a) **Expliciter la vraisemblance / log vraisemblance**

La probabilité d'avoir obtenu le tirage $X_{(n,1)}$ de puis la v.a. X (id $D=1$)

$$\mathbb{P}(X = X_{(n,1)}) = \mathbb{P}(x_1 = x_1, x_2 = x_2, \dots, x_N = x_N)$$

$$= \prod_{i=1}^N \prod_{j=1}^D \mathbb{P}(X_{ij} = x_{ij})$$

$$= \prod_{i=1}^N \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x_{in} - \mu)^2}{2\sigma^2}\right)$$

log-vraisemblance:

$$\mathcal{L}(\theta) = N \log \left(\frac{1}{\sqrt{2\pi\sigma^2}} \right) + \sum_{i=1}^N \log \left(\exp\left(-\frac{(x_{in} - \mu)^2}{2\sigma^2}\right) \right) = -\frac{N}{2} \log(2\pi\sigma^2) + \sum_{i=1}^N \left(-\frac{(x_{in} - \mu)^2}{2\sigma^2} \right)$$

NOTE (me): pour $X_{nd} \sim \text{Exp}(\lambda, d)$

$$\mathcal{L}(\theta) = \mathbb{P}(X = X_{(1,D)}) = \mathbb{P}(\forall n \leq N, \forall d \leq D, X_{nd} = x_{nd}) = \prod_{n=1}^N \prod_{d=1}^D \mathbb{P}(X_{nd} = x_{nd}, d)$$

$$\mathbb{P}(X_{nd} = x) = \lambda e^{-\lambda x} \text{ d'où } \mathcal{L}(\theta) = \prod_{n=1}^N \prod_{d=1}^D \lambda e^{-\lambda x_{nd}} \quad (= \prod_{n=1}^N \lambda e^{-\lambda x_n \text{ en dim } 1})$$

$$\log(\mathcal{L}(\theta)) = \sum_{n=1}^N \sum_{d=1}^D (\log \lambda - \lambda x_{nd}) \quad (= N \log \lambda - \lambda \sum_{n=1}^N x_n \text{ en dim } 1)$$

$$= N \log \lambda - \lambda \sum_{n=1}^N x_n \text{ en dim } 1$$

$$= N \log \lambda - \lambda \sum_{n=1}^N x_n \text{ en dim } 1$$

$$= N \log \lambda - \lambda \sum_{n=1}^N x_n \text{ en dim } 1$$

b) **En déduire μ_{MLE}**

$$\mu_{MLE} = \arg \max_{\mu} \mathcal{L} : 0 = \frac{\partial}{\partial \mu} \mathcal{L} \Rightarrow 0 = 0 + \frac{\partial}{\partial \mu} \sum_{n=1}^N \left(-\frac{(x_{n1} - \mu)^2}{2\sigma^2} \right)$$

$$\Rightarrow 0 = -\sum_{n=1}^N \frac{(x_{n1} - \mu)}{\sigma^2} \quad (-1) \Rightarrow 0 = \sum_{n=1}^N (x_{n1} - \mu)$$

$$\Rightarrow \mu = \frac{1}{N} \sum_{n=1}^N x_n$$

avec **D dimensions**: a) $\theta = \{?, \dots\}$

b) $\theta^* = \dots$ (avant)

c) **Quelle est la probabilité pour un point $\vec{x}_{test}$ d'appartenir à une certaine classe k ?**

(on estime la vraisemblance de la donnée)

$$\text{Bayes: } \mathbb{P}(y=k | \vec{x}, \theta) = \frac{\mathbb{P}(\vec{x} | y=k, \theta) \pi_k}{\sum_{j=1}^K \mathbb{P}(\vec{x} | y=j, \theta) \pi_j} \Rightarrow \mathbb{P}(y=k | \vec{x}, \theta) \propto \mathbb{P}(\vec{x} | y=k, \theta) \pi_k$$

$$\mathbb{P}(y=k | \vec{x}, \theta) = \frac{\pi_k \mathbb{P}(\vec{x} | y=k, \theta)}{\sum_{j=1}^K \pi_j \mathbb{P}(\vec{x} | y=j, \theta)}$$

version normalisée

$$\text{Ici: } \mathbb{P}(y=k | \vec{x}_{test}, \theta) \propto \frac{1}{\sqrt{2\pi\sigma_{k,d}^2}} \exp\left(-\frac{(x_{test,d} - \mu_{k,d})^2}{2\sigma_{k,d}^2}\right)$$

$$\propto \frac{1}{\sqrt{2\pi\sigma_{k,d}^2}} \exp\left(-\frac{(x_{test,d} - \mu_{k,d})^2}{2\sigma_{k,d}^2}\right)$$

d) i) Un modèle Gaussien Naïf n'est pas approprié car les différents attributs (features) sont a priori corrélés entre eux. Le choix du Naïf est problématique. Le choix gaussien, passablement.

ii) **Comment contourner l'inconvénient?** on peut passer au modèle gaussien non-naïf, c.à.d.

avec une matrice de covariance pleine. Mais on a alors KD paramètres pour les μ et $O(KD^2)$ pour la covariance

($\propto O(KD^2)$)

(*) **Lien entre Bayésien et K-moyennes**

$\vec{\mu}_k$ K-moyennes = les $\vec{\mu}_k$ du mélange de Gaussiennes

π_k = proba pour le point n d'être issu du cluster k $\rightarrow$ des probas ds $\mathcal{Q}(c)$

Bayésien plus général: (1) affectations molles (on donne une probabilité pour chaque cluster) 2) K-means

suppose des clusters de m taille / dimension alors qu'avec un mélange gaussien chaque cluster peut avoir son rayon

σ_k

Dernière partie du cours: TF-IDF, K-moyennes, arbres de décision binaire.

I. Notions de base

1) Arbres de décision binaire

Ideée générale: On se donne (x_i, y_i) avec $x_i \in \mathbb{R}^D$ (features) et y_i une classe

un arbre de décision binaire fait une suite de questions binaires "feature $x_j \leq t$?" et envoie chaque point dans une feuille.

Chaque feuille contient des exemples (échantillons) d'une ou plusieurs classes.

La prédiction de la feuille peut être: soit la classe majoritaire (vote à la majorité simple)

Soit un vecteur de probabilités de classe. L'espace des features est ainsi découpé en régions

rectangulaires alignées sur les axes. On note D l'ensemble des exemples arrivant dans un nœud (ou une feuille)

Measures d'impureté: Soit p_k la proportion d'exemples de classe k dans D

Taux d'erreur de classification: $\text{Err}(D) = 1 - \max_k p_k$

Indice de Gini: $\text{Gini}(D) = 1 - \sum_k p_k^2$

Entropie: $\text{Ent} = - \sum_k p_k \log p_k$

Pour une subdivision $D \rightarrow (D_{\text{gauche}}, D_{\text{droite}})$ par $(x_j \leq t)$ on définit $\Delta I(j, t) = \text{Ent}(D) - \frac{|D_{\text{gauche}}|}{|D|} \text{Ent}(D_{\text{gauche}}) - \frac{|D_{\text{droite}}|}{|D|} \text{Ent}(D_{\text{droite}})$

On choisit (j, t) qui maximise ΔI (i.e. minimise l'impureté des deux feuilles)

Algo: (choix de la coupure) Pour chaque feature

· Pour chaque seuil possible

· Calculer la pureté de la division possible

· Prendre la meilleure combinaison (meilleure feature avec son meilleur seuil)

· Diviser les données et selon cette coupure

options d'arbres: Vote majoritaire (si un nœud contient plusieurs classes, on peut en faire une feuille et prédire la

classe majoritaire / descente jusqu'à pureté (on continue à diviser tant que la feuille contient plusieurs classes)

→ d'où l'idée de choisir une

fonction adéquate pour évaluer la pureté (via validation)

→ overfitting

→ train loss

→ profondeur

→ largeur

II. TF-IDF

Définitions de base:

tokens: (occurrence) lexical (racine) → part- dans porteur, portif... / grammatical → ex: s au pluriel, -ait de l'imparfait...

- un token est une occurrence d'une unité linguistique dans un texte

type: (forme unique)

- un type est la forme unique du token

vocabulaire V (dictionnaire des types)

- liste des types distincts rencontrés

taille d'un texte = nombre de tokens du texte

richesse de vocabulaire d'un texte: nombre de types apparaissant

(taille de l'ensemble V des types présents)

- On va compter 1 lemme (par exemple: verbe aimer) comme ~70 types

a) Problèmes avec différentes langues / où couper les mots → 70 formes flexionnelles en français)

Exemples:

l'allemand est composé de mots extrêmement longs car souvent composé d'autres mots avec beaucoup de déclinaisons

- pomme de terre pourra être pris comme trois mots séparés, d'abord on se signification s'il était pris en un mot

- le chinois et le japonais n'ont pas d'alphabets de 'mots' mais plutôt de caractères

- ponctuation, apostrophes, mots composés (choix de découpage)

- langues avec morphologie complexe

(un seul mot peut contenir plusieurs morphèmes)

pb: explosion du nombre de types si pas de lemmatisation

ex: كرتي "un seul mot pour dire 'il l'a écrit'"

- langues avec agglutination (coréen, hongrois): les mots s'entraînent en blocs → segmentation indispensable

problèmes des homographes:

2 mots ayant la même écriture mais dont la signification est différente

mots ambigus: mots s'écrivant de la même manière mais dont la signification est totalement différente (ex: avocat)

polysémies: " " sign. différente mais se rejoignent (ex: bureau)

idée générale:

idée: comparer les fréquences des types entre les documents

→ classer des documents par thème devrait être possible car certains types sont spécifiques à un thème

L'idée générale est de faire en sorte que les termes les plus courants importent moins (le résultat du calcul, aient un point faible dans le calcul du vecteur représentant un texte)

term. Frequency: $TF(n, d) = \# \text{ fois où le type } d \text{ apparaît dans le doc } n$ (→ c'est en fait un comptage)

Document Frequency: $DF(d) = \# \text{ docs où le type } d \text{ apparaît au moins une fois}$ types très courants:

Inverse document Frequency: IDF

$= \log \left(\frac{N_{\text{docs}}}{DF(d)} \right)$ version régulière

$IDF(d) = \log \left(\frac{1 + N_{\text{docs}}}{1 + DF(d)} \right) + 1$ types très rares:

la formule TF-IDF est le produit $TF(d) \cdot IDF(d)$ (aport d'un doc relatif avec tous les types)

nb de @ après classificateur linéaire: V pour une classification binaire

KV pour multiclass si Bayésien avec un seul @ par exemple / one vs rest

one vs one: $V \frac{K(K-1)}{2}$ paramètres.

matrice de corrélation:

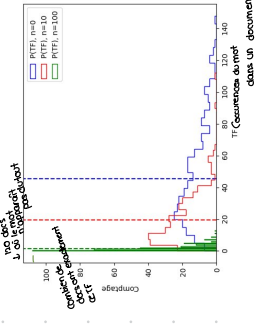
[comment lire?]

→ chaque case (i, j) = similitude entre l'exemple i et l'exemple j (valeur entre -1 et 1)

→ couleur claire → corrélation forte → exemples semblables

→ couleurs foncées → corrélation faible → exemples différents

la diagonale est toujours très claire ($\text{corr}(x, x) = 1$) → on l'ignore



i.e. la loi de Zipf dit que, dans un grand corpus, la fréquence d'un mot est inversement proportionnelle à son rang: si un mot est le r -ième plus fréquent, $f(r) \approx C/r$. Quelques mots apparaissent énormément, et la plupart des mots apparaissent très rarement.

