

FONDEMENTS DU MACHINE LEARNING

1^{er} octobre 2024

Aujourd'hui :

Cours (13^h45 - 15^h15) : ACP

TD (15^h30 - 18^h45) : Fin exercices SVD

TD/TP ACP

Analyse en composantes principales (ACP) pour la réduction de dimension

Problème: $X \in \mathbb{R}^{m \times n}$ matrice de données

$$X = \begin{bmatrix} x_1^T \\ \vdots \\ x_m^T \end{bmatrix}, \quad x_i \in \mathbb{R}^n \text{ vecteur représentatif de l'individu } i, \text{ à } n \text{ attributs}$$

But: Réduire X à une matrice de taille $m \times k$ avec $k \ll n$ en conservant le maximum d'information pour chacun des individus $x_1, \dots, x_m$

$\Rightarrow$ Analyse en k composantes principales

① Analyse en 1 composante principale

$\hookrightarrow$ Objectif: Remplacer les n attributs par une nouvelle variable que l'on appelle composante principale.

$\Rightarrow$ Géométriquement, cela revient à représenter le nuage de points $\{x_1, \dots, x_m\} \subseteq \mathbb{R}^n$ par $\{y_1, \dots, y_m\} \subseteq \mathbb{R}^n$ tels que les $\{y_i\}$ appartiennent à une même droite

Deux points de vue: Géométrie: On reste dans $\mathbb{R}^n$ et on passe d'un nuage de points quelconque à un nuage de points sur une droite

Algèbre $X \in \mathbb{R}^{m \times n} \rightarrow c \in \mathbb{R}^{m \times 1}$

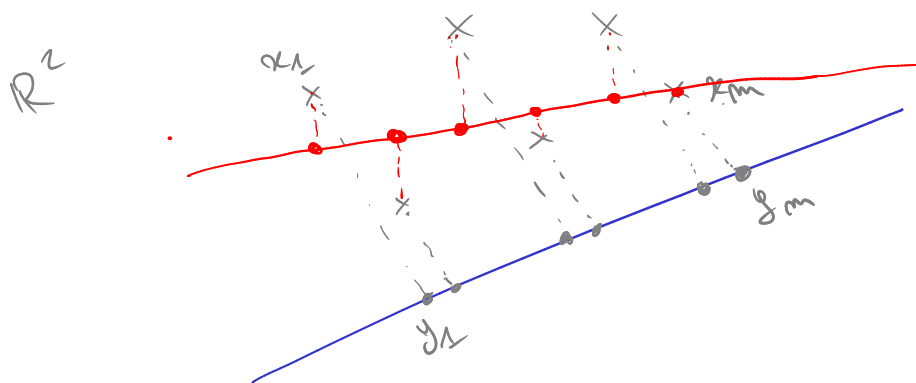
$\Rightarrow$ La meilleure façon de réduire les m attributs en 1 caractéristique
 à déterminer la meilleure projection des $\{x_i\}$ sur une droite
 dans $\mathbb{R}^m$

$\Rightarrow$ La "meilleure droite possible" (au sens de fournir la représentation
 des données la plus proche des points de départ) est celle qui maximise
 l'inertie du nuage de points projetés

$X = \{x_1, \dots, x_m\}$: nuage de points de départ

Inclue $Y = \{y_1, \dots, y_m\}$: nuage des points projetés, contenus dans
 une droite / ligne.

$$\left(\text{Inertie de } Y : \frac{1}{m-1} \sum_{i=1}^m \|y_i - \bar{y}\|^2 \quad \bar{y} = \frac{1}{m} \sum_{i=1}^m y_i \in \mathbb{R}^m \right)$$



$\hookrightarrow$ Comment calcule-t-on les y_i / la droite en pratique ?

• Une "bonne" droite devrait passer par l'individu moyen
 $\bar{x} = \frac{1}{m} \sum_{i=1}^m x_i$ afin de conserver l'information moyenne du nuage

de points $\Rightarrow$ Les coordonnées des projections du nuage de
 points ne prennent pas en compte la variance
 centrale des données (mais les y_i, s_i !)

• Les $\{y_i\}$ appartiennent à une droite d'équation

$$\{y = \bar{x} + \gamma u \mid \gamma \in \mathbb{R}\}$$

où u est le vecteur directeur de la droite à déterminer

On a alors : $y_i = \bar{x} + c_i u$ avec $c_i \in \mathbb{R}$

↑
composante principale de x_i

$$X \in \mathbb{R}^{m \times n} \rightarrow C \in \mathbb{R}^{m \times 1}, \quad C = \begin{bmatrix} c_1 \\ \vdots \\ c_m \end{bmatrix} \text{ vecteurs des premières composantes principales}$$

Théorème :

Pour tout $x \in \mathbb{R}^n$, la projection de x sur la droite $\{y = \bar{x} + \gamma u\}$ est donnée par

$$\bar{x} + u^T (x - \bar{x}) u$$

Corollaire : Les projections $\{y_1, \dots, y_m\}$ du nuage de points

$\{x_1, \dots, x_m\}$ sont données par

$$y_i = \bar{x} + u^T (x_i - \bar{x}) u \quad \forall i=1..m$$

↳ On cherche le vecteur u tel que l'inertie du nuage de points projetés $\mathcal{Y} = \{y_1, \dots, y_m\}$ soit maximale

$$I(\mathcal{Y}) = \frac{1}{m-1} \sum_{i=1}^m \|y_i - \bar{y}\|^2$$

$$= \frac{1}{m-1} \sum_{i=1}^m \|\bar{x} + u^T (x_i - \bar{x}) u - \bar{y}\|^2$$

$$\text{On } \bar{y} = \frac{1}{m} \sum_{i=1}^m y_i$$

$$= \frac{1}{m} \sum_{i=1}^m (\bar{x} + u^T (x_i - \bar{x}) u)$$

$$= \left(\frac{1}{m} \sum_{i=1}^m \bar{x} \right) + \left(\frac{1}{m} \sum_{i=1}^m u^T (x_i - \bar{x}) u \right)$$

$$= \bar{x} + \frac{1}{m} u^T \left(\sum_{i=1}^m (x_i - \bar{x}) \right) u$$

$$\begin{aligned} \sum_{i=1}^m u^T (x_i - \bar{x}) u &= \left[\sum_{i=1}^m u^T (x_i - \bar{x}) \right] u \\ &= \left[u^T \left(\sum_{i=1}^m (x_i - \bar{x}) \right) \right] u \end{aligned}$$

Donc

$$I(y) = \frac{1}{m-1} \sum_{i=1}^m \left\| \bar{x} + u^T (x_i - \bar{x}) u - \bar{y} \right\|^2$$

$$= \frac{1}{m-1} \sum_{i=1}^m \left\| \bar{x} + u^T (x_i - \bar{x}) u - \bar{x} - \frac{1}{m} u^T \sum_{i=1}^m (x_i - \bar{x}) u \right\|^2$$

$$= \frac{1}{m-1} \sum_{i=1}^m \left\| u^T (x_i - \bar{x}) u - \frac{1}{m} u^T \sum_{i=1}^m (x_i - \bar{x}) u \right\|^2$$

$$= \frac{1}{m-1} \sum_{i=1}^m \left\| u^T \left[(x_i - \bar{x}) - \frac{1}{m} \sum_{i=1}^m (x_i - \bar{x}) \right] u \right\|^2$$

$$\begin{aligned} x_i - \bar{x} - \frac{1}{m} \sum_{i=1}^m (x_i - \bar{x}) &= x_i - \bar{x} - \frac{1}{m} \sum_{i=1}^m x_i + \bar{x} = x_i - \frac{1}{m} \sum_{i=1}^m x_i \\ &= x_i - \bar{x} \end{aligned}$$

$$I(y) = \frac{1}{m-1} \sum_{i=1}^m \left\| u^T (x_i - \bar{x}) u \right\|^2$$

$$\|v\|^2 = v^T v$$

$$= \frac{1}{m-1} \sum_{i=1}^m u^T (x_i - \bar{x})^T u u^T (x_i - \bar{x}) u$$

$$= \frac{1}{m-1} \sum_{i=1}^m \left[(x_i - \bar{x})^T u \right]^2 u^T u$$

$$= \frac{1}{m-1} \sum_{i=1}^m \left[u^T (x_i - \bar{x}) \right]^2 \|u\|^2$$

$$= \frac{1}{m-1} \sum_{i=1}^m u^T (x_i - \bar{x}) (x_i - \bar{x})^T u \times \|u\|^2$$

$$(a^T b)^2 = (a^T b)(a^T b) = (a^T b)(b^T a) = a^T (b b^T) a$$

$$I(u) = \frac{1}{m-1} \|u\|^2 \sum_{i=1}^m u^T (x_i - \bar{x}) (x_i - \bar{x})^T u$$

$$= \|u\|^2 u^T \left(\frac{1}{m-1} \sum_{i=1}^m (x_i - \bar{x}) (x_i - \bar{x})^T \right) u$$

$$= \|u\|^2 u^T \Sigma u$$

↑
 Σ : matrice de covariance
 de $\{x_1, \dots, x_m\}$

↳ Pour maximiser cette matrice par rapport à u , on a 2 cas :

- Soit $x_i = \bar{x} \ \forall i$ et dans ce cas $u = 0$
 $\hookrightarrow \Sigma = [0]$
- Sinon, on prend u de norme 1 et égal à un vecteur propre de Σ associé à la plus grande valeur propre de Σ

Formule à retenir

ACP avec 1 composante principale

$$\{x_1, \dots, x_m\} \subseteq \mathbb{R}^n \longrightarrow \{y_1, \dots, y_m\} \subseteq \mathbb{R}^n$$

avec $y_i = \bar{x} + u^T (x_i - \bar{x}) u$, $\forall i=1..m$, u vecteur propre de norme 1 associé à la plus grande valeur propre de $\Sigma = \frac{1}{m-1} \sum_{i=1}^m (x_i - \bar{x}) (x_i - \bar{x})^T$

c_i : (1^{re}) composante principale de x
 u : (1^{er}) vecteur principal
 $\{\bar{x} + \gamma u\}$: (1^{er}) axe principal

Toutes ces notions sont déterminées par le nuage de points X

En pratique, on calcule les points projetés via

$$\underbrace{Y}_{m \times m} = \begin{bmatrix} y_1^T \\ \vdots \\ y_m^T \end{bmatrix} = \underbrace{\left(\underbrace{X}_{m \times n} - \underbrace{1_m}_{m \times 1} \underbrace{\bar{x}^T}_{1 \times n} \right)}_{m \times n} \underbrace{u}_{n \times 1} \underbrace{u^T}_{1 \times n} + \underbrace{1_m}_{m \times 1} \underbrace{\bar{x}^T}_{1 \times n}$$

(2) Analyse en k composantes principales, $k \geq 1$

But: $X \in \mathbb{R}^{m \times n} \rightarrow C = \begin{bmatrix} c_1^T \\ \vdots \\ c_m^T \end{bmatrix} \in \mathbb{R}^{m \times k}$

$$X = \{x_1, \dots, x_m\} \in \mathbb{R}^n \rightarrow Y = \{y_1, \dots, y_m\} \in \mathbb{R}^m$$

avec $\{y_i\} \subseteq \{ \bar{x} + U_k \gamma \mid \gamma \in \mathbb{R}^k \}$

Sous-espace
affine de
dimension k

et $U_k = [u_1 \dots u_k] \in \mathbb{R}^{n \times k}$ formé par k vecteurs linéairement indépendants et de norme 1

$\Rightarrow$ On écrit: $y_i = \bar{x} + U_k c_i$

$c_i \in \mathbb{R}^k$ est le vecteur des composantes principales

↳ Le meilleur choix possible pour les U_k consiste à prendre $u_1, \dots, u_k$ orthogonaux ($u_i^T u_j = 0 \text{ } i \neq j$) et à les prendre comme k vecteurs propres de Σ associés aux k plus grandes valeurs propres

NB: On exclut le cas $\Sigma = [0]$, pour lequel on peut prendre $U_k = [0]$

↳ On a alors $y_i = \bar{x} + U_k C_i$, avec $C_i = \begin{bmatrix} u_1^T (x_i - \bar{x}) \\ \vdots \\ u_k^T (x_i - \bar{x}) \end{bmatrix}$

C_i : vecteur des k premières composantes principales de x_i

U_k : k premiers vecteurs principaux

$\{\bar{x} + U_k \delta \mid \delta \in \mathbb{R}^k\}$: sous-espace associé aux k premiers axes principaux

$$\{\bar{x} + \delta u_1 \mid \delta \in \mathbb{R}\}, \{\bar{x} + \delta u_2 \mid \delta \in \mathbb{R}\}, \dots, \{\bar{x} + \delta u_k \mid \delta \in \mathbb{R}\}$$

Propriété importante: Calculer les composantes principales peut se faire en calculant les composantes principales une à une (notamment parce Σ possède une base orthogonale de vecteurs propres)

(3) Le cas $k=n$

$$X \in \mathbb{R}^{m \times n} \longrightarrow C \in \mathbb{R}^{m \times n}$$

$$x_i \in \mathbb{R}^n \longrightarrow y_i = \bar{x} + U_n C_i \quad C_i \in \mathbb{R}^n$$

$[u_1, \dots, u_n]$ base orthogonale de vecteurs propres de Σ

c_i : coordonnées de $x_i - \bar{x}$ dans la base $(u_1, \dots, u_m)$

Avec $k=m$, on fait un changement de base (+ une translation par rapport à $\bar{x}$) relativement à notre jeu de données $x_1, \dots, x_m$

$\Rightarrow$ Dans cette base, les coordonnées des $x_i - \bar{x}$ sont décorrélées

Représentation naturelle de X	Représentation avec l'ACP C
Base canonique $\{e_1, \dots, e_n\}$	Base de vecteurs propres de Σ
$X^c = X - I_m \bar{x}^T$ (Données centrées)	$C^c = \underbrace{X^c}_{n \times m} \underbrace{U_m}_{m \times m}$
$\Sigma = \frac{(X^c)^T X^c}{m-1}$	$\Sigma' = \frac{(C^c)^T (C^c)}{m-1} = U_m^T \Sigma U_m$ $= \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_m \end{bmatrix}$ où $\lambda_1 \geq \dots \geq \lambda_m$ sont les valeurs propres de Σ $\Rightarrow$ Représentation dans laquelle les m attributs sont décorrélés