

Réduction de dimensionnalité

Sergey Kirgizov

Algèbre linéaire et problèmes d'optimisation

- Matrices

- Multiplication

- Norme, produit scalaire

- Valeurs et vecteurs propres

- Quotient de Rayleigh

- Théorème min-max de Courant-Fischer

Analyse en composantes principales

- Matrice de données

- Algorithmes de calcul

Nous utilisons :

- ▶ des lettres majuscules en gras (**A**, **B**, **C**, ...) pour désigner des matrices
- ▶ des lettres minuscules en gras (**a**, **b**, **c**, ...) pour les vecteurs
- ▶ des lettres minuscules (a , b , c , $\dots$, α , β , γ , $\dots$, λ , $\dots$) pour les scalaires.

$$\mathbf{v} = \begin{pmatrix} v_1 \\ v_2 \\ v_3 \\ \vdots \\ v_n \end{pmatrix}$$

$$\mathbf{v}^T = \begin{pmatrix} v_1 & v_2 & v_3 & \dots & v_n \end{pmatrix}$$

Matrices

$$A_{n \times k} = \begin{pmatrix} a_{1,1} & a_{1,2} & \cdots & a_{1,k} \\ a_{2,1} & a_{2,2} & \cdots & a_{2,k} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n,1} & a_{n,2} & \cdots & a_{n,k} \end{pmatrix} = A$$

La dimension de A est (n, k)

Matrices

$$A_{n \times k} = \begin{pmatrix} a_{1,1} & a_{1,2} & \cdots & a_{1,k} \\ a_{2,1} & a_{2,2} & \cdots & a_{2,k} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n,1} & a_{n,2} & \cdots & a_{n,k} \end{pmatrix} = A$$

La dimension de A est (n, k)

Operations

Transposition $C = A^T \Leftrightarrow c_{j,i} = a_{j,i}$

Addition $A_{n \times k} = B_{n \times k} + C_{n \times k} \Leftrightarrow a_{i,j} = b_{i,j} + c_{i,j}$

Multiplication par un scalaire $B = cA \Leftrightarrow b_{i,j} = ca_{i,j}$

Multiplication $C_{n \times m} = A_{n \times k} B_{k \times m} \Leftrightarrow c_{i,j} = \sum_{p=1}^k a_{i,p} b_{p,j}$

Attention aux dimensions !

Multiplication

Matrice ligne **L** par une matrice colonne **C**

$$\mathbf{LC} = \begin{pmatrix} x_1 & \dots & x_k \end{pmatrix} \begin{pmatrix} y_1 \\ \vdots \\ y_k \end{pmatrix} = \sum_{i=1}^k x_i y_i$$

Multiplication

Matrice ligne **L** par une matrice colonne **C**

$$\mathbf{LC} = \begin{pmatrix} x_1 & \dots & x_k \end{pmatrix} \begin{pmatrix} y_1 \\ \vdots \\ y_k \end{pmatrix} = \sum_{i=1}^k x_i y_i$$

Cas général

L_i : les lignes de la première matrice

C_j : les colonnes de la deuxième

$$\begin{pmatrix} \boxed{L_1} \\ \vdots \\ \boxed{L_m} \end{pmatrix} \begin{pmatrix} \boxed{C_1} & \dots & \boxed{C_n} \end{pmatrix} = \begin{pmatrix} L_1 C_1 & L_1 C_2 & \dots & L_1 C_n \\ L_2 C_1 & L_2 C_2 & \dots & L_2 C_n \\ \vdots & \vdots & \ddots & \vdots \\ L_m C_1 & L_m C_2 & \dots & L_m C_n \end{pmatrix}$$

Algorithmes et leurs complexité

Ici nous supposons que l'addition et la multiplications prennent $O(1)$ secondes. C'est vrai en pratique pour les nombres codés avec un nombre d'octets limité. Sinon, il faut voir les algorithmes de Karatsuba, Toom-Cook, Schönhage-Strassen, Fürer.

$$\mathbf{A}_{n \times k} \mathbf{B}_{k \times m}$$

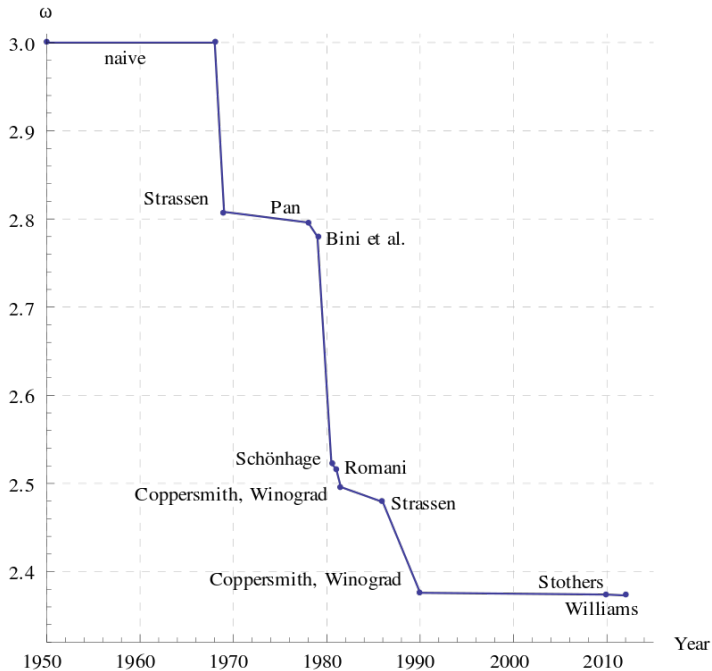
- ▶ **Algorithme classique** $O(nkm)$.

Dans le cas de matrices carrées $O(n^3)$

- ▶ **Strassen algorithm** $O(N^{\log_2 7}) \approx O(N^{2.8074})$

<http://www.cs.utsa.edu/~wagner/CS3343/strassen/strassen.html>

$$O(n^\omega)$$



Bah ! ça reste compliqué...

Même la lecture des
valeurs prend $O(n^2)$

Ordinateurs quantiques ??

Multiplication rapide pour le big data matrice ?

Deux solutions :

- ▶ Algorithmes parallèles et distribués
- ▶ Utiliser (si possible) les matrices creuses :
les matrices contenant beaucoup de zéros.
Idée : ne pas stocker les zéros. Ne rien multiplier par 0, le résultat est connu d'avance !

MLlib: Main Guide

- [Basic statistics](#)
- [Data sources](#)
- [Pipelines](#)
- [Extracting, transforming and selecting features](#)
- [Classification and Regression](#)
- [Clustering](#)
- [Collaborative filtering](#)
- [Frequent Pattern Mining](#)
- [Model selection and tuning](#)

Data Types - RDD-based API

- [Local vector](#)
- [Labeled point](#)
- [Local matrix](#)
- [Distributed matrix](#)
 - [RowMatrix](#)
 - [IndexedRowMatrix](#)
 - [CoordinateMatrix](#)
 - [BlockMatrix](#)

MLlib supports local vectors and matrices stored on a single machine, as well as distributed matrices backed by one or more RDDs. Local vectors and local matrices are simple data models that serve as public interfaces. The underlying linear algebra operations are provided by [Breeze](#). A training example used in supervised learning is called a "labeled point" in MLlib.

https:
[//spark.apache.org/docs/latest/mllib-data-types.html](https://spark.apache.org/docs/latest/mllib-data-types.html)

MLlib: Main Guide

- [Basic statistics](#)
- [Data sources](#)
- [Pipelines](#)
- [Extracting, transforming and selecting features](#)
- [Classification and Regression](#)
- [Clustering](#)
- [Collaborative filtering](#)
- [Frequent Pattern Mining](#)
- [Model selection and tuning](#)
- [Advanced topics](#)

MLlib: RDD-based API Guide

- [Data types](#)
- [Basic statistics](#)
- [Classification and regression](#)
- [Collaborative filtering](#)
- [Clustering](#)
- [Dimensionality reduction](#)
- [Feature extraction and transformation](#)
- [Frequent pattern mining](#)
- [Evaluation metrics](#)
- [PMML model export](#)

Distributed matrix

A distributed matrix has long-typed row and column indices and double-typed values, stored distributively in one or more RDDs. It is very important to choose the right format to store large and distributed matrices. Converting a distributed matrix to a different format may require a global shuffle, which is quite expensive. Four types of distributed matrices have been implemented so far.

The basic type is called `RowMatrix`. A `RowMatrix` is a row-oriented distributed matrix without meaningful row indices, e.g., a collection of feature vectors. It is backed by an RDD of its rows, where each row is a local vector. We assume that the number of columns is not huge for a `RowMatrix` so that a single local vector can be reasonably communicated to the driver and can also be stored / operated on using a single node. An `IndexedRowMatrix` is similar to a `RowMatrix` but with row indices, which can be used for identifying rows and executing joins. A `CoordinateMatrix` is a distributed matrix stored in [coordinate list \(COO\)](#) format, backed by an RDD of its entries. A `BlockMatrix` is a distributed matrix backed by an RDD of `MatrixBlock` which is a tuple of `(Int, Int, Matrix)`.

Note

The underlying RDDs of a distributed matrix must be deterministic, because we cache the matrix size. In general, the use of non-deterministic RDDs can lead to errors.

RowMatrix

A `RowMatrix` is a row-oriented distributed matrix without meaningful row indices, backed by an RDD of its rows, where each row is a local vector. Since each row is represented by a local vector, the number of columns is limited by the integer range but it should be much smaller in practice.

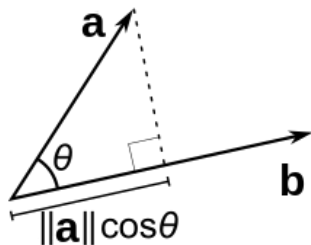
[Scala](#)
[Java](#)
[Python](#)

A `RowMatrix` can be created from an `RDD[Vector]` instance. Then we can compute its column summary statistics and decompositions. [QR decomposition](#) is of the form $A = QR$ where Q is an orthogonal matrix and R is an upper triangular matrix. For [singular value decomposition \(SVD\)](#) and [principal component analysis \(PCA\)](#), please refer to [Dimensionality](#)

Norme de vecteur, produit scalaire

Vecteurs $\mathbf{a} = \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{pmatrix}$ et $\mathbf{b} = \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{pmatrix}$ de dimension n .

$$\langle \mathbf{a}, \mathbf{b} \rangle = \mathbf{a}^T \mathbf{b} = \|\mathbf{a}\| \|\mathbf{b}\| \cos \theta = \sum_{i=1}^n a_i b_i$$



$$\|\mathbf{a}\|^2 = \langle \mathbf{a}, \mathbf{a} \rangle$$

$$\mathbf{A}\mathbf{x} = \lambda\mathbf{x}$$

$$\mathbf{A}\mathbf{x} = \lambda\mathbf{x}$$

- ▶ λ valeur propre
- ▶ $\mathbf{x}$ vecteur propre

Les vissage propres (Eigenfaces).

An individual face was video recorded and digitized at $2^7 \times 2^7$ pixels with 2^8 gray level by means of an IVS-100 image processor. Faces were lined up by a cross-hair overlay display that appeared on a video monitor. The vertical line passed through the symmetry line of the face and the horizontal line through the pupils of the eyes. Field depth was adjusted so that facial width was the same for images. Since these steps were all adjusted by eye, this contributed to the general error level. The pictures were taken under background-lighting conditions. Since the lighting varied with the time of day, this too was a source of error. To some extent, this error was diminished by a normalization procedure that we now describe.

A face, or for that matter any object, can be regarded as a pointwise map of reflectivities, say, $r(\mathbf{x})$. Under a uniform illumination, say, I , the face is given by

$$\tilde{\varphi}(\mathbf{x}) = Ir(\mathbf{x}). \quad (21)$$

For a variety of reasons, it is important to normalize a picture so that a reference portion of a face is at a standard level of illumination. If we denote a reference point by $\mathbf{x}_0$ and standard light level by I_0 , then we take the normalized picture to be

$$\varphi(\mathbf{x}) = \frac{I_0}{\varphi(\mathbf{x}_0)} \tilde{\varphi}(\mathbf{x}). \quad (22)$$

In actual practice we took the reference portion to be small, high cheek areas below each eye and averaged the light level over these two patches. This procedure provides a specific light-level normalization. In addition, it provides the basis for the future identification of a picture.

The average face was computed according to Eq. (3) and is

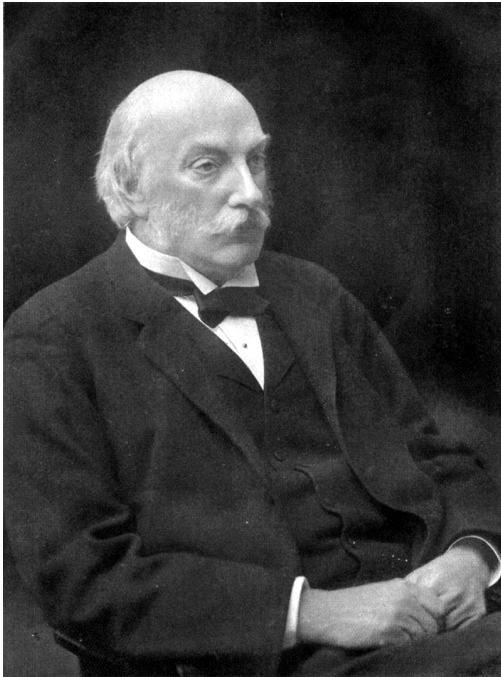
Sirovich and Kirby (1987)
Turk and Pentland (1991)



Caractérisation variationnelle des valeurs propres !

Minima d'une certaine
fonction !

The Lord Rayleigh /'reɪli/



Quotient de Rayleigh

Une matrice symétrique $A_{n,n}$, un vecteur $x = (x_1, x_2, \dots, x_n)^T$

$$R(A, x) = \frac{x^T A x}{x^T x}$$

$$R(A, x) = \frac{\sum_{i=1}^n a_{i,j} x_i x_j}{\sum_{i=1}^n x_i^2}$$

On s'intéresse qu'aux vecteurs de la norme 1, c'est-à-dire

$$\sum_{i=1}^n x_i^2 = 1, \text{ alors } R(A, x) = \sum_{i,j=1}^n a_{i,j} x_i x_j$$

Théorème min-max de Courant-Fischer

Si la matrice $\mathbf{M}$ est symétrique, ces valeurs propres sont réelles :

$$\lambda_1 \geq \lambda_2 \geq \cdots \geq \lambda_n$$

Théorème min-max de Courant-Fischer

Si la matrice $\mathbf{M}$ est symétrique, ces valeurs propres sont réelles :

$$\lambda_1 \geq \lambda_2 \geq \cdots \geq \lambda_n$$

$$\lambda_1 = \max_{\|\mathbf{x}\|=1} \left(\mathbf{x}^T \mathbf{A} \mathbf{x} \right)$$

$$\rho_1 = \arg \max_{\|\mathbf{x}\|=1} \left(\mathbf{x}^T \mathbf{A} \mathbf{x} \right)$$

$$\mathbf{A} \mathbf{v}_1 = \lambda_1 \mathbf{v}_1, \|\mathbf{v}_1\| = 1$$

Théorème min-max de Courant-Fischer

Si la matrice $\mathbf{M}$ est symétrique, ces valeurs propres sont réelles :

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$$

$$\lambda_1 = \max_{\|\mathbf{x}\|=1} \left(\mathbf{x}^T \mathbf{A} \mathbf{x} \right)$$

$$\mathbf{p}_1 = \arg \max_{\|\mathbf{x}\|=1} \left(\mathbf{x}^T \mathbf{A} \mathbf{x} \right)$$

$$\mathbf{A} \mathbf{v}_1 = \lambda_1 \mathbf{v}_1, \|\mathbf{v}_1\| = 1$$

$$\lambda_2 = \max_{\|\mathbf{x}\|=1, \mathbf{x} \perp \mathbf{p}_1} \left(\mathbf{x}^T \mathbf{A} \mathbf{x} \right)$$

$$\mathbf{p}_2 = \arg \max_{\|\mathbf{x}\|=1, \mathbf{x} \perp \mathbf{p}_1} \left(\mathbf{x}^T \mathbf{A} \mathbf{x} \right)$$

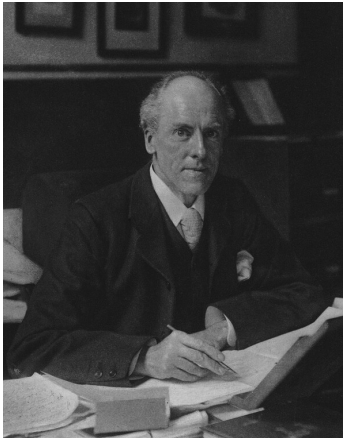
$$\mathbf{A} \mathbf{p}_2 = \lambda_2 \mathbf{p}_2, \|\mathbf{p}_2\| = 1$$

etc...

Analyse en composantes principales

*PCA was invented in 1901 by Karl Pearson, as an analogue of the **principal axis theorem** in mechanics; it was later independently developed and named by Harold Hotelling in the 1930s.*

— Wikipedia



Matrice de données

Matrice de données

$$\mathbf{D}_{n \times k} = \begin{pmatrix} d_{1,1} & d_{1,2} & \cdots & d_{1,k} \\ d_{2,1} & d_{2,2} & \cdots & d_{2,k} \\ \vdots & \vdots & \ddots & \vdots \\ d_{n,1} & d_{n,2} & \cdots & d_{n,k} \end{pmatrix}$$

n nombre des instances de données.

k nombre de caractéristiques (features).

PCA. Minimisation des erreurs

Matrice de données

$$D_{n \times k} = \begin{pmatrix} d_{1,1} & d_{1,2} & \cdots & d_{1,k} \\ d_{2,1} & d_{2,2} & \cdots & d_{2,k} \\ \vdots & \vdots & \ddots & \vdots \\ d_{n,1} & d_{n,2} & \cdots & d_{n,k} \end{pmatrix} = \begin{pmatrix} \boxed{\delta_1} \\ \boxed{\delta_2} \\ \vdots \\ \boxed{\delta_n} \end{pmatrix}$$

n nombre des instances de données

k nombre de caractéristiques (features)

δ_i ième instance.

On cherche un vecteur x , $\|x\| = 1$, t.q.

$$\min \sum_{i=1}^n \|\delta_i - \langle \delta_i, x \rangle x\|^2$$

On minimise la somme des différences entre les données et leurs projections sur x !

$$\langle \delta_i, x \rangle = \|\delta_i\| \|x\| \cos \theta = \|\delta_i\| \cos \theta$$

$\langle \delta_i, x \rangle \in \mathbb{R}$, mais $\langle \delta_i, x \rangle x$ est un vecteur !

Minimisation des erreurs = Maximisation de variance !

On cherche un vecteur $\mathbf{x}$, $\|\mathbf{x}\| = 1$, t.q.

$$\min \sum_{i=1}^n \|\delta_i - \langle \delta_i, \mathbf{x} \rangle \mathbf{x}\|^2$$

Minimisation des erreurs = Maximisation de variance !

On cherche un vecteur $\mathbf{x}$, $\|\mathbf{x}\| = 1$, t.q.

$$\min \sum_{i=1}^n \|\delta_i - \langle \delta_i, \mathbf{x} \rangle \mathbf{x}\|^2$$

Théorème de Pythagore

$$\|\delta_i - \langle \delta_i, \mathbf{x} \rangle \mathbf{x}\|^2 = \|\delta_i\|^2 - \langle \delta_i, \mathbf{x} \rangle^2$$

Minimisation des erreurs = Maximisation de variance !

On cherche un vecteur $\mathbf{x}$, $\|\mathbf{x}\| = 1$, t.q.

$$\min \sum_{i=1}^n \|\delta_i - \langle \delta_i, \mathbf{x} \rangle \mathbf{x}\|^2$$

Théorème de Pythagore

$$\|\delta_i - \langle \delta_i, \mathbf{x} \rangle \mathbf{x}\|^2 = \|\delta_i\|^2 - \langle \delta_i, \mathbf{x} \rangle^2$$

Alors,

$$\max \sum_{i=1}^n \langle \delta_i, \mathbf{x} \rangle^2$$

Maximisation de variance

On cherche un vecteur $\mathbf{x}$, $\|\mathbf{x}\| = 1$, t.q.

$$\max \sum_{i=1}^n \langle \delta_i, \mathbf{x} \rangle^2 = \max \|\mathbf{D}\mathbf{x}\|^2 =$$

$$\|\mathbf{a}\|^2 = \langle \mathbf{a}, \mathbf{a} \rangle$$

$$= \max \langle \mathbf{D}\mathbf{x}, \mathbf{D}\mathbf{x} \rangle$$

$$\langle \mathbf{a}, \mathbf{b} \rangle = \mathbf{a}^T \mathbf{b}$$

$$= \max ((\mathbf{D}\mathbf{x})^T \mathbf{D}\mathbf{x})$$

$$(\mathbf{AB})^T = \mathbf{B}^T \mathbf{A}^T$$

$$= \max (\mathbf{x}^T \mathbf{D}^T \mathbf{D}\mathbf{x})$$

$$\text{Associativité } (\mathbf{AB})\mathbf{C} = \mathbf{A}(\mathbf{BC})$$

$$= \max \left(\mathbf{x}^T (\mathbf{D}^T \mathbf{D}) \mathbf{x} \right)$$

$\mathbb{E}[Z]$ désigne l'espérance mathématique de la variable aléatoire Z . Pour deux variables aléatoires réelles X et Y on définit la **covariance**

$$\text{Cov}(X, Y) = \mathbb{E} \left[(X - \mathbb{E}[X]) (Y - \mathbb{E}[Y]) \right]$$

La **variance** de X est donc $\text{Var}(X) = \text{Cov}(X, X)$.

L'**écart type** de X est $\sigma(X) = \sqrt{\mathbb{E} \left[(X - \mathbb{E}[X])^2 \right]}$.

Matrice de n-covariance

Considérons chaque colonne comme un échantillon d'une variable aléatoire.

Matrice de n-covariance

Considérons chaque colonne comme un échantillon d'une variable aléatoire.

$$\mathbf{D}_{n \times k} = \begin{pmatrix} \mathbf{v}_1 & \mathbf{v}_2 & \cdots & \mathbf{v}_k \end{pmatrix}$$

Matrice de n-covariance

Considérons chaque colonne comme un échantillon d'une variable aléatoire.

$$\mathbf{D}_{n \times k} = \begin{pmatrix} \mathbf{V}_1 & \mathbf{V}_2 & \cdots & \mathbf{V}_k \end{pmatrix}$$

La matrice de n-covariance est la matrice $\mathbf{C}$, t.q $c_{i,j}$ est égal à la covariance entre $\mathbf{V}_i$ et $\mathbf{V}_j$ fois n .

$$\mathbf{B} = \begin{pmatrix} \mathbf{V}_1 - \mathbb{E}[\mathbf{V}_1] & \mathbf{V}_2 - \mathbb{E}[\mathbf{V}_2] & \cdots & \mathbf{V}_k - \mathbb{E}[\mathbf{V}_k] \end{pmatrix}$$

Matrice

$$\mathbf{C} = \mathbf{B}^T \mathbf{B}$$

est symétrique semi-définie positive, ces valeurs propres sont réelles et positives :

$$\lambda_1 \geq \lambda_2 \geq \cdots \geq 0$$

Application du Théorème min-max de Courant-Fischer

On cherche un vecteur $\mathbf{x}$, $\|\mathbf{x}\| = 1$, t.q.

$$\max \left(\mathbf{x}^T (\mathbf{D}^T \mathbf{D}) \mathbf{x} \right)$$

Si on change la matrice $\mathbf{D}$ en soustrayant la moyenne de chaque colonne... on pourra interpreter $\mathbf{D}^T \mathbf{D}$ comme la matrice de n-covariance $\mathbf{C}$.

Application du Théorème min-max de Courant-Fischer

On cherche un vecteur $\mathbf{x}$, $\|\mathbf{x}\| = 1$, t.q.

$$\max \left(\mathbf{x}^T (\mathbf{D}^T \mathbf{D}) \mathbf{x} \right)$$

Si on change la matrice $\mathbf{D}$ en soustrayant la moyenne de chaque colonne... on pourra interpreter $\mathbf{D}^T \mathbf{D}$ comme la matrice de n-covariance $\mathbf{C}$.

$$\max \left(\mathbf{x}^T \mathbf{C} \mathbf{x} \right)$$

Application du Théorème min-max de Courant-Fischer

On cherche un vecteur $\mathbf{x}$, $\|\mathbf{x}\| = 1$, t.q.

$$\max \left(\mathbf{x}^T (\mathbf{D}^T \mathbf{D}) \mathbf{x} \right)$$

Si on change la matrice $\mathbf{D}$ en soustrayant la moyenne de chaque colonne... on pourra interpreter $\mathbf{D}^T \mathbf{D}$ comme la matrice de n-covariance $\mathbf{C}$.

$$\max \left(\mathbf{x}^T \mathbf{C} \mathbf{x} \right)$$

et par le théorème min-max de Courant-Fischer on sait que c'est la plus grande valeur propre qui maximise ce quotient de Rayleigh.

Schéma d'algorithme PCA

- ▶ Soustraire des moyennes des colonnes
- ▶ Calculer la matrice de n-covariance C
- ▶ Trouver les plus grandes valeurs et vecteurs propres de la matrice C
- ▶ Projetez les données sur le sous-espace engendré par ces vecteurs. Plus il y a des vecteurs, mieux on explique les données.

On n'a pas besoin de calculer toutes les vecteurs propres.

Nous ne sommes intéressés que par ceux qui sont liés aux plus grandes valeurs propres.

Et encore... on peut les approcher sans calculer avec une précision absolue par des méthodes itératives :

- ▶ Méthode de la puissance itérée
- ▶ Itération du quotient de Rayleigh

MLlib: Main Guide

- [Basic statistics](#)
- [Data sources](#)
- [Pipelines](#)
- [Extracting, transforming and selecting features](#)
- [Classification and Regression](#)
- [Clustering](#)
- [Collaborative filtering](#)
- [Frequent Pattern Mining](#)
- [Model selection and tuning](#)
- [Advanced topics](#)

MLlib: RDD-based API Guide

- [Data types](#)
- [Basic statistics](#)
- [Classification and regression](#)
- [Collaborative filtering](#)
- [Clustering](#)
- [Dimensionality reduction](#)
 - [singular value decomposition \(SVD\)](#)
 - [principal component analysis \(PCA\)](#)
 - [Feature extraction and](#)

Principal component analysis (PCA)

Principal component analysis (PCA) is a statistical method to find a rotation such that the first coordinate has the largest variance possible, and each succeeding coordinate, in turn, has the largest variance possible. The columns of the rotation matrix are called principal components. PCA is used widely in dimensionality reduction.

« spark.mllib supports PCA for tall-and-skinny matrices stored in row-oriented format and any Vectors.

[Scala](#)
[Java](#)
[Python](#)

The following code demonstrates how to compute principal components on a `RowMatrix` and use them to project the vectors into a low-dimensional space.

Refer to the [RowMatrix Python docs](#) for details on the API.

```
from pyspark.mllib.linalg import Vectors
from pyspark.mllib.linalg.distributed import RowMatrix

rows = sc.parallelize([
    Vectors.sparse(5, {1: 1.0, 3: 7.0}),
    Vectors.dense(2.0, 0.0, 3.0, 4.0, 5.0),
    Vectors.dense(4.0, 0.0, 0.0, 6.0, 7.0)
])

mat = RowMatrix(rows)
# Compute the top 4 principal components.
# Principal components are stored in a local dense matrix.
pc = mat.computePrincipalComponents(4)

# Project the rows to the linear space spanned by the top 4 principal components.
projected = mat.multiply(pc)
```

Autres méthodes de réduction de dimensionnalité :

- ▶ Non-negative matrix factorization (NMF)
- ▶ Kernel PCA
- ▶ Graph-based kernel PCA
- ▶ Linear discriminant analysis (LDA)
- ▶ Generalized discriminant analysis (GDA)
- ▶ Autoencoder

Questions ?