

知的情報処理 11-12. 仲間を集めるク ラスタリング(1,2)

櫻井彰人
慶應義塾大学理工学部

クラスタとは

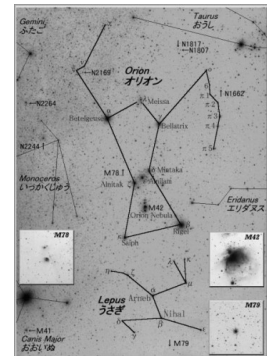
- クラスタとは cluster のこと。
- cluster とは房のこと。
 - 葡萄の房
- 小さいものが複数個集まって、大きな一個の塊を構成するとき、この塊をクラスタという
 - あ、忌まわしきクラスタ爆弾の「クラスタ」もこれ



クラスタリング・クラスタ分析とは？

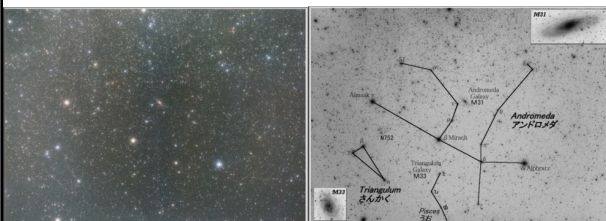
- クラスタ: 対象(データ)の集合
 - 同一クラスタ内では、似ている
 - 他のクラスタの対象とは似ていない
- クラスタ分析
 - 対象の集合を(複数の)クラスタに集める
- クラスタリングには、普通、正解がない
 - 葡萄の房はつながっているか否かでわかるが、「似ている」か否かの判断には、多くの視点があるし、程度問題でもあるから
- 正解があっても、与えられていない or 人知には知れない
- どんなときに使うか
 - スタンドアロンとして: 対象(データ)の分布に関する知見を得るため
 - 他のアルゴリズムの前処理として

星座: クラスタリングか？



<http://homepage3.nifty.com/chokainomori-ao/gallery002.htm>

星座



<http://homepage3.nifty.com/chokainomori-ao/gallery002.htm>

クラスタリングの評価: 均質性と分離性

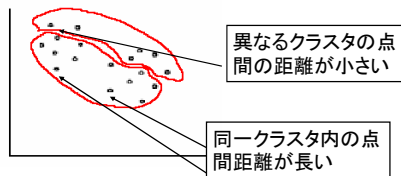
- ・ **均質性:** 同一クラスタ内の要素は互いに類似している
- ・ **分離性:** 異なるクラスタの要素は互いに異なっている
- ・ ...クラスタリングは容易ではない

こうしたデータが与えられたとき、クラスタリングアルゴリズムは2つのクラスタを見出す。次のように。



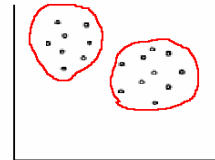
悪いクラスタリング

このクラスタリングは、均質性も分離性も満たしていない



よいクラスタリング

このクラスタリングなら、均質性と分離性を満たしている



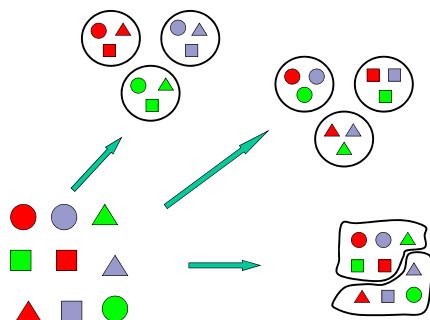
よいクラスタリングとは？

- よいクラスタ分析手法はよいクラスターを作る(はず)。よいクラスターとは
 - クラスター内の類似性は高い
 - クラスター間の類似性は低い
- クラスタリングの結果のよさは、用いる類似性尺度と手法のよさによる。
- クラスタリング手法のよさは、「隠れた」パターン(規則性)が発見できる能力で測ることができよう。
- なお、クラスタリングというのは、主観的であることに注意されたい
 - 正解がない。データが不足していて正解が決定できない場合だけでなく、本来に正解がない場合(見方によって変わる場合)がほとんど

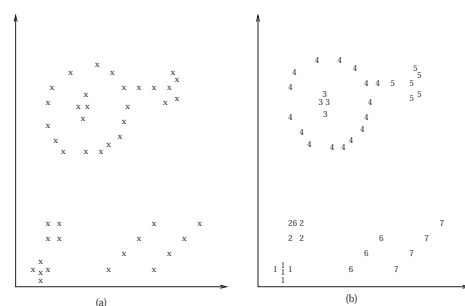
クラスタリングに対する要件

- スケーラビリティがある
- 異なった属性が扱える
- どんな形のクラスターでも発見できる
- パラメータを決めるために必要とされる領域知識の少ない
- ノイズや外れ値 (outlier) があってもよい
- データの入力順序に依存しない
- 次元が高くても大丈夫
- ユーザが制約を与えることができる
- 他の手法と組み合わせることができる

そうはいっても難しい



意地悪な例ならいくらでも考えられる



意地悪ではない: 軸のスケールを変えると

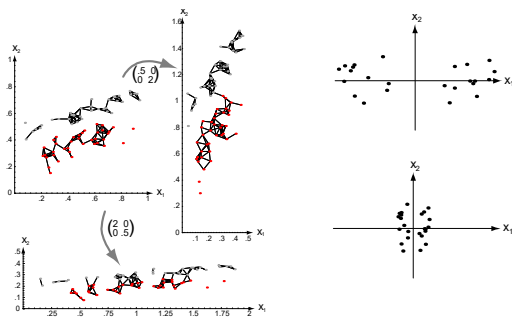


FIGURE 10.8. From: Richard O. Duda, Peter, E. Hart, and David G. Stork. Pattern Classification. Copyright c 2001 by John Wiley & Sons, Inc.

FIGURE 10.9. From: Richard O. Duda, Peter, E. Hart, and David G. Stork. Pattern Classification. Copyright c 2001 by John Wiley & Sons, Inc.

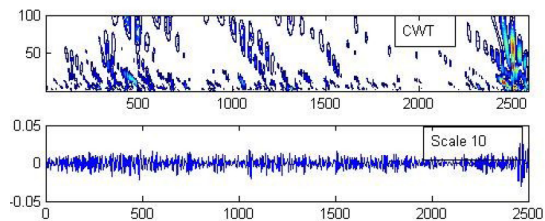
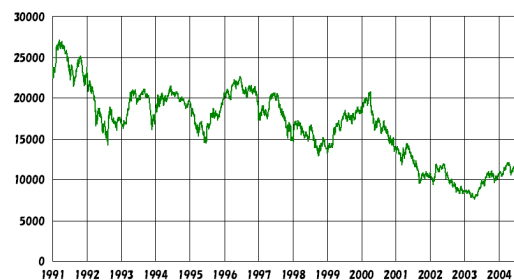
属性(特徴)選択とスケーリング

- 適切な属性(特徴)とスケーリングを選ぶことが必要
- 適切な属性(特徴)を選ぶことは難しい
 - クラスタリングの正解がわかっているならば、いろいろ工夫ができる。
 - しかし、そうでなければ(正解がわかっていないのが普通)コンピュータにはできない、というのが正しい
 - その分野の専門家しかできないだろう。
- スケーリング、より正確には、距離の定義が重要
 - 適切な距離が分かるということは、しかし、適切なクラスタリングが分かるということ
 - 従って、これも、難しい

別の重要な(or 根源的な)問題

- そもそも、クラスタというのは存在するのか？
- 人間に「クラスタ」「仲間」「繰り返し性」が見えれば、クラスタが存在すると考えていいのか？
- 見えなければいけないのか？
 - 高次元空間のクラスタは見えないぞ

クラスタ(パターン)はあるか？



株価: クラスタはあるか？



株式市場に現れる謎の“羅針盤” 騰落率の放射模様は株価予測に役立つ?
<http://www.nikkei.co.jp/needs/analysis/07/a070906.html>

株価: クラスタはあるか？ (もと)

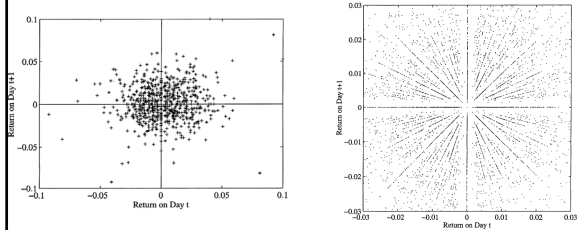


Figure 1. From Braley and Myers-Weyerhaeuser. This is a reproduction of Figure 13-2 from Braley and Myers (1991). Compare this to Figure 2 below. The horizontal axis shows return on Weyerhaeuser stock on a given day; the vertical axis shows return on the next day. Pluses (+) are used to symbolize the data. The data span the period January 2, 1980 to December 30, 1988. Out of 758 points, 6 lie outside the bounds of the plot.

Crack, Timothy F., and Olivier Ledoit. 1996. "Robust Structure Without Predictability: The 'Compass Rose' Pattern of the Stock Market." *Journal of Finance*, 51: 751-762.

株価: クラスタはあるか？ (もと)

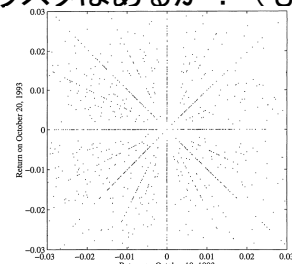


Figure 3. Compass Rose-2123 NYSE stocks. This plot shows one point only for each of 2123 NYSE stocks on a randomly selected pair of consecutive trading days; the structure is cross-sectionally coherent. The horizontal axis shows return on October 19, 1993; the vertical axis shows return on October 20, 1993. Out of 2123 points, 469 lie outside the bounds of the plot.

Crack, Timothy F., and Olivier Ledoit. 1996. "Robust Structure Without Predictability: The 'Compass Rose' Pattern of the Stock Market." *Journal of Finance*, 51: 751-762.

Henrik Amilon (henrik.amilon@nek.lu.se) and Hans Byström
The Compass Rose Pattern of the Stock Market: How Does it Affect Parameter Estimates, Forecasts, and Statistical Tests?
Working Papers, Department of Economics, Lund University, No 2000-18

音楽: クラスタリングによるある結果

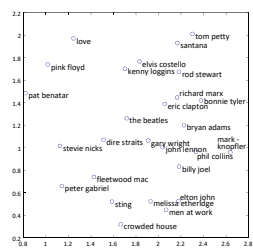
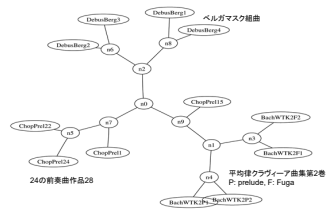


Figure 1: Artists embedded in a 2-D space. This is a small portion of the full space derived from the Er dds measure.

D. Ellis, B. Whitman, A. Berenzweig, S. Lawrence.
The Quest for Ground Truth in Musical Artist Similarity
Proc. ISMIR-02, Paris, October 2002.



R. Cilibrasi, P.M.B. Vitányi, R. de Wolf, Algorithmic clustering of music based on string compression, *Computer Music J.*, 28:4 (2004), 49-67.

なお、ショパンの前奏曲Op. 28は、バウハの平均値クワイア曲集に収められたと言われている

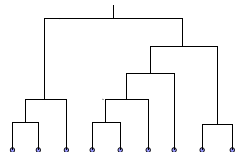
クラスタリング技法

- 階層的方法 hierarchical: クラスタ間の関係が樹木状になるように構成する方法。根は全体。葉は(通常は)一要素からなるクラスタ。枝分かれが、クラスタの分割(逆にみるとクラスタの融合)に相当する
 - 凝集法 agglomerative: 各要素がそれぞれ一つのクラスタを構成する状態から開始。近いクラスタを順々にまとめて大きなクラスタとしていく
 - 分割法 divisive: 全体を一つのクラスタとして開始。より小さなクラスタに分割していく
- 非階層的方法 non- hierarchical: 階層的でないもの
 - モデル法: 統計的モデル(クラスター分布)を考え、尤度最大化
 - 実際のアルゴリズムとしては、EMアルゴリズムが著名
 - Mathematica の ClusterAnalysis はこれに属する
 - 分割法: 分割の良否の評価関数の最適化問題と考える
 - k-means法
 - グラフ理論に基づく方法、spectral clustering等

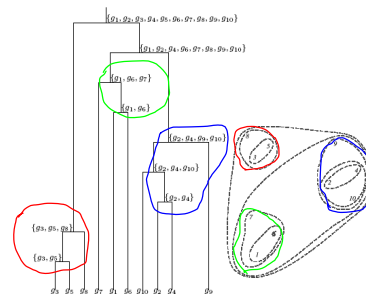
注: Dendrogram 樹形図

対象(データ)集合の融合(逆向きに見れば分割)の様子を2進木の形に表したものだ。ただし、例えば下図であれば、縦軸が融合時の非類似度を表すようにする。

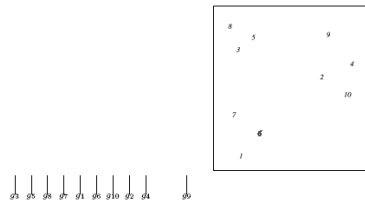
必要と考えるクラスタ数の枝があるところで、切断すれば、任意数クラスタへのクラスタリングが得られる。



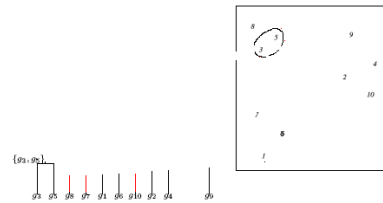
階層的クラスタリング



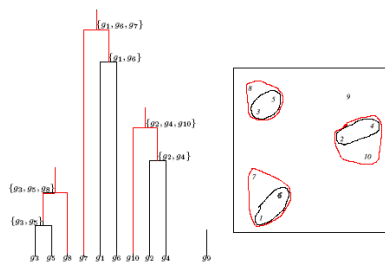
階層的クラスタリング: 凝集法



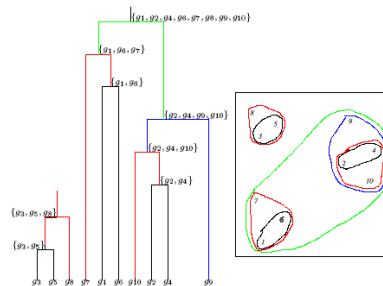
階層的クラスタリング: 凝集法



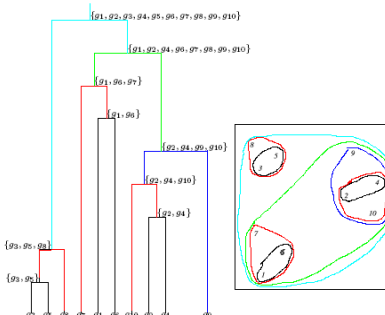
階層的クラスタリング: 凝集法



階層的クラスタリング: 凝集法

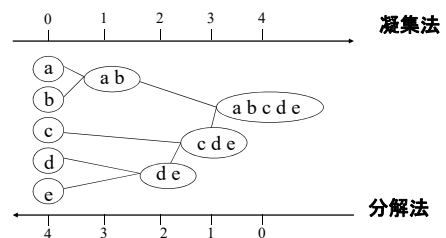


階層的クラスタリング: 凝集法



階層的クラスタリング

- 距離行列をクラスタリングの基準に用いる。この方法では、クラスタ数 k を入力パラメータとしない。



凝集法: アルゴリズム

本アルゴリズムは、点の個数 n と $n \times n$ の距離行列 (各点間の距離を要素とする。対称行列となる) を入力とする

1. 凝集法 (d, n)
2. 各要素からなるクラスタを作る。 n 個とする
3. 各クラスタが一つの頂点 (葉) であるような木 T を作る
4. while 2個以上のクラスタがある
5. 最も近い二つのクラスタを見つける。 C_i と C_j とする
6. C_i と C_j を一緒にして新クラスタ C を作る。要素数 $|C_i| + |C_j|$
7. **C から他のすべてのクラスタへの距離を計算する**
8. T に新たな頂点 C を付加し、頂点 C_i 及び C_j に結ぶ
9. クラスタ間の距離を保持する表 d から C_i と C_j に関する情報を削除する
10. 新クラスタ C と他のクラスタとの距離を、表 d に付け加える
11. T を値として終了

距離の定義が異なれば、異なるクラスタ・クラスタ構造が得られる

クラスタ間距離

- 最小 (最近) 距離:

$$d_{\min}(C_i, C_j) = \min_{p \in C_i, p' \in C_j} d(p, p')$$

- 最大 (最遠) 距離:

$$d_{\max}(C_i, C_j) = \max_{p \in C_i, p' \in C_j} d(p, p')$$

- 平均距離:

$$d_{\text{average}}(C_i, C_j) = \frac{1}{n_i n_j} \sum_{p \in C_i} \sum_{p' \in C_j} d(p, p')$$

- 中心 (重心) 距離:

$$d_{\text{mean}}(C_i, C_j) = d(m_i, m_j)$$

- Ward 距離

$$d_{\text{Ward}}(C_i, C_j) = e(C_i \cup C_j) - e(C_i) - e(C_j)$$

$$e(C) = \sum_{x \in C} (d(x, c))^2 \quad (c \text{ は } C \text{ の中心})$$

クラスタ間の距離: 再掲

- $d_{\min}(C, C^*) = \min_{\text{全要素 } x \in C \text{ と } y \in C^*} d(x, y)$
 - 二つのクラス間距離は、すべての要素対の距離の最小値
- $d_{\max}(C, C^*) = \max_{\text{全要素 } x \in C \text{ と } y \in C^*} d(x, y)$
 - 二つのクラス間距離は、すべての要素対の距離の最大値
- $d_{\text{avg}}(C, C^*) = (1 / |C^*| |C|) \sum_{\text{全要素 } x \in C \text{ と } y \in C^*} d(x, y)$
 - 二つのクラス間距離は、すべての要素対の距離の平均値

距離と類似度と非類似度

- 類似度 (similarity): 大きい = 距離が近い
 - ちょっと変わった例: 文字列の編集距離 (edit distance)
 - 例: 文字列の Jaro-Winkler distance
- 非類似度 (dissimilarity): 小さい = 距離が近い

点間の距離

- 数学でポピュラーなのはミンコフスキー距離:

$$d(i, j) = \sqrt[q]{(|x_{i1} - x_{j1}|^q + |x_{i2} - x_{j2}|^q + \dots + |x_{ip} - x_{jp}|^q)}$$

ただし $i = (x_{i1}, x_{i2}, \dots, x_{ip})$ と $j = (x_{j1}, x_{j2}, \dots, x_{jp})$ は p -次元データ、 q はある正数 (通常は整数)

- 特別な場合 (コンピュータではこっちがポピュラー): $q = 1$ のとき、マンハッタン距離と呼ばれる

$$d(i, j) = |x_{i1} - x_{j1}| + |x_{i2} - x_{j2}| + \dots + |x_{ip} - x_{jp}|$$

点間の距離

- $q = 2$ ならばユークリッド距離:

$$d(i, j) = \sqrt{(|x_{i1} - x_{j1}|^2 + |x_{i2} - x_{j2}|^2 + \dots + |x_{ip} - x_{jp}|^2)}$$

- 距離の定義

- $d(i, j) \geq 0$
- $d(i, i) = 0$
- $d(i, j) = d(j, i)$
- $d(i, j) \leq d(i, k) + d(k, j)$

- 他には、荷重付きの距離、Pearson の積率相関係数等々

補足: 距離らしくない距離

- Compression distance (圧縮距離) というものもある
- データを圧縮するアルゴリズムに基づく
 - 例えば、テキストAを、テキストB(の統計情報)を知って圧縮すると、知らないで圧縮するより短くできたとする。そうすると、テキストAとテキストBは似ている、すなわち、距離はかなり近いと考えられる。
 - これは、数学的な奥行きが非常に深い概念である

補足: 編集距離

- 2つの文字列の近さを計る方法
 - かな漢字変換エンジンや、DNAの相同性検索などに利用される
- 2つの文字列の 編集距離
 - 片方の文字列から、もう一方の文字列を得るための操作回数の最小値
 - 使用可能な操作は、例えば、以下の3種類
 - 削除: 1つの文字を取り除く
 - 挿入: 1つの文字を新たに付け加える
 - 置換: 1つの文字を別の文字で置き換える
 - 参考になるURI「wikipedia:レーベンシュタイン距離」と「GPU Challenge 2009」
<http://ja.wikipedia.org/wiki/%E3%83%AC%E3%83%BC%E3%83%99%E3%83%B3%E3%82%B7%E3%83%A5%E3%82%BF%E3%82%A4%E3%83%B3%E8%B7%9D%E9%9B%A2>
<http://sacsis.hpcc.jp/2009/gpu/toolkit10.ppt>

階層的クラスタリング: 使う「距離」

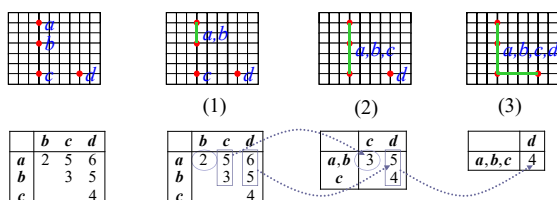
- single-link クラスタリング (または connectedness/minimum 法): 二つのクラスタ間の距離を、要素間の最小距離で定義する。類似性という言葉で表現するなら、二つのクラスタの類似性は、最もよく似た要素対間の類似性であると定義することになる(ちょっと苦しいぞ)。
- complete-link クラスタリング (または the diameter/maximum 法): 二つのクラスタ間の距離を、要素間の最大距離で定義する。類似性という言葉で表現するなら、二つのクラスタの類似性は、最も異なった要素対間の類似性であると定義することになる(ちょっと苦しいぞ)。
- average-link クラスタリング: 二つのクラスタ間の距離を、全要素間の距離の平均値で定義する。

クラスタ間の距離: 補足

- single-link クラスタリング (または connectedness/minimum 法): 二つのクラスタ間の距離を、要素間の最小距離で定義する。類似性という言葉で表現するなら、二つのクラスタの類似性は、最もよく似た要素対間の類似性であると定義することになる(ちょっと苦しいぞ)。
 - complete-link クラスタリング (または the diameter/maximum 法): 二つのクラスタ間の距離を、要素間の最大距離で定義する。類似性という言葉で表現するなら、二つのクラスタの類似性は、最も異なった要素対間の類似性であると定義することになる(ちょっと苦しいぞ)。
 - average-link クラスタリング: 二つのクラスタ間の距離を、全要素間の距離の平均値で定義する。
- ☐ Single-Link / 最近隣
☐ Complete-Link / 最遠隣(?)
☐ Average 平均.
☐ Centroids 中心.

Single-Link クラスタリング

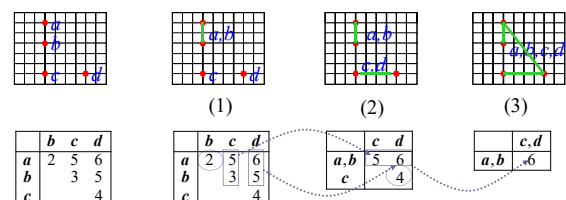
ユークリッド距離



距離行列

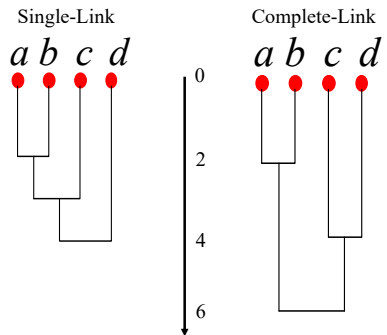
Complete-Link クラスタリング

ユークリッド距離



距離行列

デンドログラム(樹形図)で比較



距離データの構造

■ 非対称距離

$$\begin{bmatrix} x_{11} & \dots & x_{1f} & \dots & x_{1p} \\ \dots & \dots & \dots & \dots & \dots \\ x_{i1} & \dots & x_{if} & \dots & x_{ip} \\ \dots & \dots & \dots & \dots & \dots \\ x_{n1} & \dots & x_{nf} & \dots & x_{np} \end{bmatrix}$$

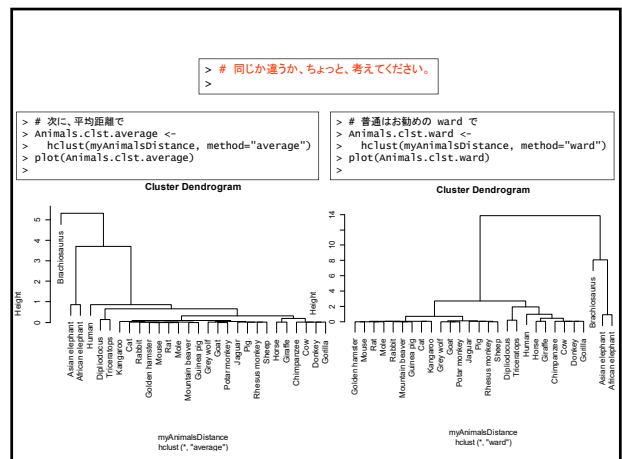
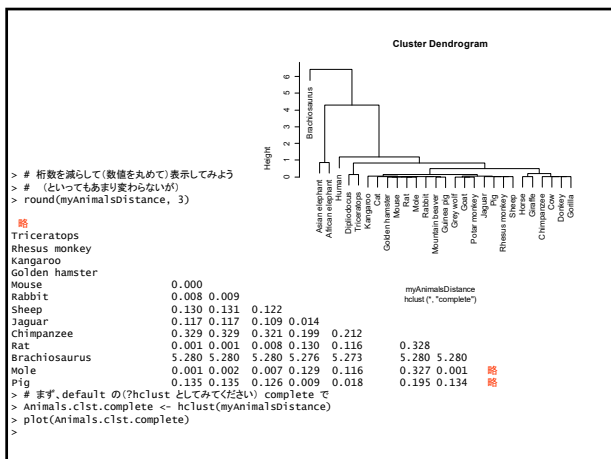
■ 対称距離

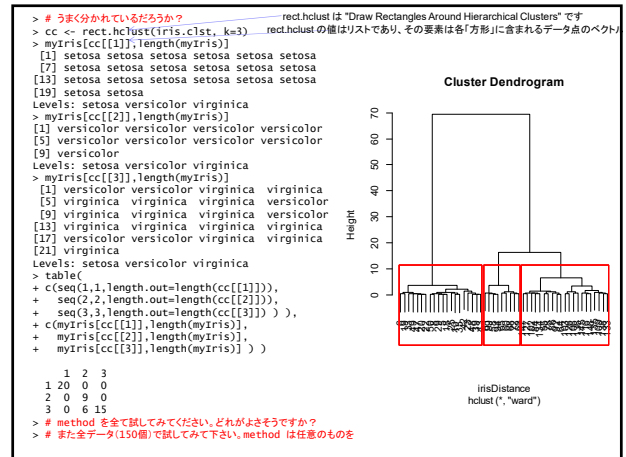
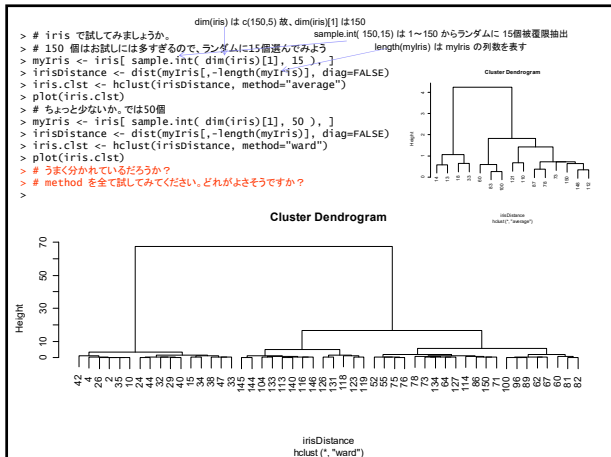
$$\begin{bmatrix} 0 & & & & \\ d(2,1) & 0 & & & \\ d(3,1) & d(3,2) & 0 & & \\ \vdots & \vdots & \vdots & \ddots & \\ d(n,1) & d(n,2) & \dots & \dots & 0 \end{bmatrix}$$

Rにおける clustering

- パッケージとしては, stats (標準に含まれている), e1071, cluster などに含まれる
- ここでは, stats の hclust (階層的)を試用する

```
> setwd("D:/R/Sample")
> Animals <- read.csv("12Animals.csv", header=TRUE)
>
> # まずデータの中味を見てみよう
> summary(Animals)
      X
African elephant: 1   Min.   : 0.023   Min.   : 0.40
Asian elephant  : 1   1st Qu.:  3.100   1st Qu.: 22.23
Brachiosaurus   : 1   Median: 53.830   Median: 137.00
Cat              : 1   Mean   : 4278.439   Mean   : 574.52
Chimpanzee      : 1   3rd Qu.: 479.000   3rd Qu.: 420.00
Cow              : 1   Max.   :87000.000   Max.   :5712.00
(Other)         :22
> # 動物名が一つの列になっている。データとしてはよくある形だが、Rとしては扱いにくい。
> # 動物名は行の名前にしよう
> myAnimals <- Animals[,2:3]
> rownames(myAnimals) <- Animals$X
> # 読み込むときに次のようにしてもよい
> myAnimals <- read.csv("12Animals.csv", header=TRUE, row.names=1)
> # 次の問題は、属性間(といっても属性しかないが)のスケールの違いである。
> # 平均0, 分散1に正規化しよう
> myAnimals <- scale(myAnimals)
> # 次も可: myAnimals <- scale(Animals[,-1]); rownames(myAnimals) <- Animals$X
>
> # クラスタリングは距離行列を元に行われる。
> # そこで、距離行列を作ろう
> # ユークリッド距離が default. 対角行は 0 なので、作らないことにしよう
> myAnimalsDistance <- dist(myAnimals, diag=FALSE)
```





本日の課題1

■ 階層的クラスタリングのスライド中にある

```

> # 同じか違うか、ちょっと、考えてください。
>

```

```

> # method を全て試してみてください。どれがよさそうですか？
> # また全データ(150個)で試してみてください。method は任意のものを

```

に答えてください。

知的情報処理

13. 仲間を集めるクラスタリング(2/2)

櫻井彰人

慶應義塾大学理工学部

k-means クラスタリング: 定式化

- 入力: n 個の点からなる集合 V , パラメータ k
- 出力: k 個の点 (クラスタの中心) からなる集合 X で、すべての可能な X のなかで、二乗歪 (squared error distortion) $d(V, X)$ が最小のもの

二乗歪

- 点 v と点集合 X が与えられたとき、 v から X への距離

$$d(v, X)$$

は、 v から X の最近接点へのユークリッド距離であるとする。

- n 個の点からなる集合 $V = \{v_1, \dots, v_n\}$ と点集合 X が与えられたとき、二乗歪 Squared Error Distortion は次のように定義される

$$d(V, X) = \sum_{1 \leq i \leq n} d(v_i, X)^2 / n$$

1-Means クラスタリング: やさしい場合

- **Input:** n 個の点からなる集合 V
- **Output:** ある一個の点 x (クラスタ中心) で、 x のすべての可能な選び方のなかで、二乗歪 $d(V, x)$ が最小となるもの

1-Means クラスタリングはやさしい問題.

しかし、中心 (クラスタ数) が2個以上となると、とたんに難しい問題 (NP-完全) となる.

効率のよい発見的方法の一つに Lloyd アルゴリズムがある

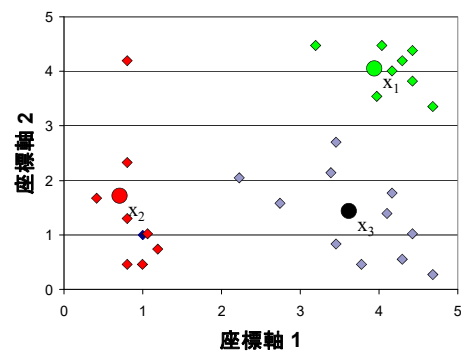
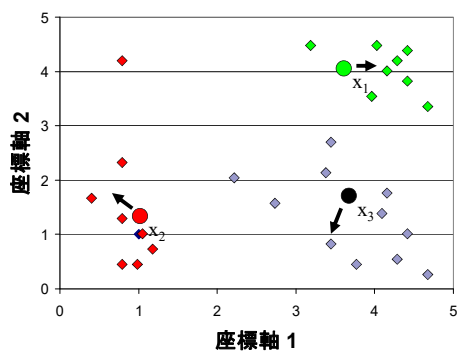
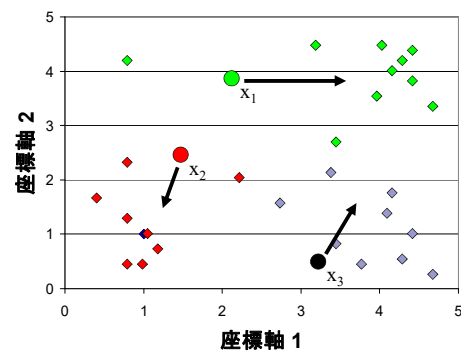
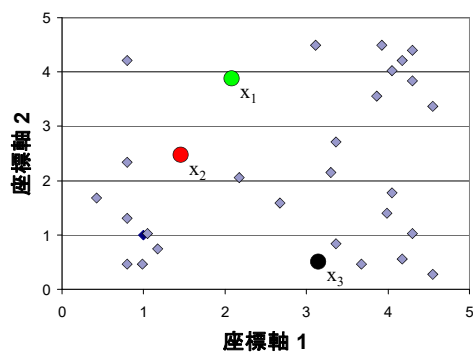
K-Means クラスタリング: Lloyd アルゴリズム

1. Lloyd アルゴリズム
2. 任意に k 個のクラスタ中心を選ぶ
3. **while** クラスタ中心が変更された
4. 各データ点をクラスタ C_i に割り当てる
割り当てるクラスタは、最も近いクラスタ中心のクラスタ ($1 \leq i \leq k$)
5. すべての点への割り当てが終わったら、
各クラスタの重心を新たなクラスタ中心とする
すなわち、新たなクラスタ中心は、クラスタ C それぞれに

$$\frac{\sum_{v \in C} v}{|C|}$$

*このアルゴリズムは局所解に陥ることがある.

このアルゴリズムは振動して停止しないことがある



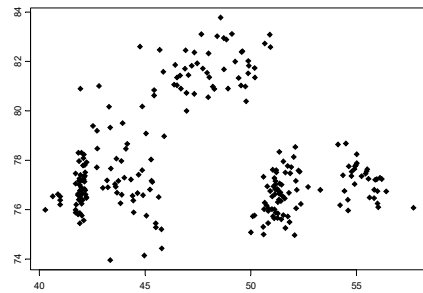
クラスタリング demo

http://home.dei.polimi.it/matteucc/Clustering/tutorial_html/AppletKM.html

<http://www.cs.washington.edu/research/imagedatabase/demo/kmcluster/>

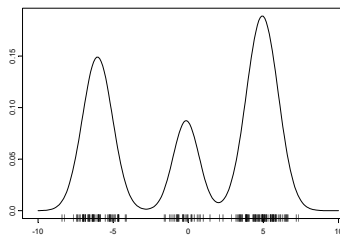
<http://tech.nitoyon.com/ja/blog/2009/04/09/kmeans-visualise/>

統計的な枠組み



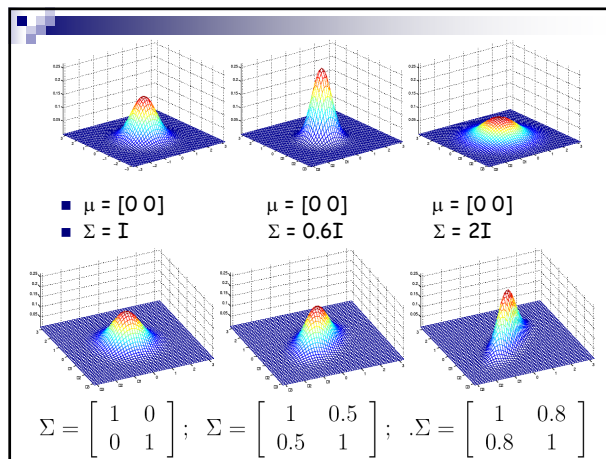
5または6個のクラスタがありそう(?)
各クラスタを正規分布でモデル化しよう
複数個の取り扱いはどうする?

ガウス混合分布:

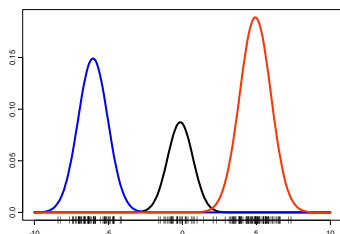


前提をおく:

- 観測値はある混合分布からのサンプルである
- $p(x) = \sum \pi_g p_g(x)$
- 各クラスタは混合分布を構成する個々の分布に対応する



ガウス混合分布



フィッティング:

π_g と $p_g(x)$ のパラメータ(平均、分散等)を推定する

尤度最大化

ガウス混合分布(ガウス分布の重み付き線形和)をフィッティングする

各 x_i はいずれかの分布(クラスタ)に属すると考える

- 対数尤度を最大化する $\pi_g, p_g()$ を求める。

$$\log L(\pi_1, \dots, \pi_G, p_1, \dots, p_G | x_1, \dots, x_n) = \log \prod_{i=1}^n \sum_{g=1}^G \pi_g p_g(x_i)$$

$$= \sum_{i=1}^n \log \left(\sum_{g=1}^G \pi_g p_g(x_i) \right)$$

尤度最大化の困難点

- 変更: 対数尤度を最大化する $\pi_g, p(\cdot|\theta_g)$ を求める。

$$\log L(\pi_1, \dots, \pi_G, \theta_1, \dots, \theta_G | x_1, \dots, x_n) = \log \prod_{i=1}^n \sum_{g=1}^G \pi_g p(x_i; \theta_g)$$

$$= \sum_{i=1}^n \log \left(\sum_{g=1}^G \pi_g p(x_i; \theta_g) \right)$$

- 困難点(等式制約もあるが):

$$\frac{\partial \log L(\pi, \theta, | \mathbf{x})}{\partial \theta_1} = \sum_{i=1}^n \frac{\pi_1 \frac{\partial p(x_i; \theta_1)}{\partial \theta_1}}{\sum_{g=1}^G \pi_g p(x_i; \theta_g)}$$

$$1 = \sum_{g=1}^G \pi_g$$

i 毎に異なる。通分すると大変なことに

しかし繰返し計算なら

$$\theta_{ML} = \arg \max_{\mu_g, \pi_g} LL(\mu_1, \dots, \pi_1, \dots)$$

$$LL(\mu_1, \dots, \pi_1, \dots) = \sum_{i=1}^n \log \sum_{g=1}^G \pi_g N(x_i; \mu_g)$$

とりあえず、停留点が見つかるかどうか、Lagrange関数を微分してみよう

$$\text{Lagrange関数は } L = \sum_{i=1}^n \log \sum_{g=1}^G \pi_g N(x_i; \mu_g) + \lambda(1 - \sum_{g=1}^G \pi_g)$$

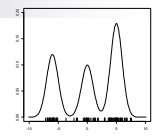
$$\frac{\partial L}{\partial \pi_g} = \sum_{i=1}^n \frac{N(x_i; \mu_g)}{\sum_{g=1}^G \pi_g N(x_i; \mu_g)} - \lambda$$

$$= \sum_{i=1}^n \tau_i^g / \pi_g - \lambda$$

$$\frac{\partial L}{\partial \mu_g} = \sum_{i=1}^n \frac{\pi_g N(x_i; \mu_g)}{\sum_{j=1}^G \pi_j N(x_i; \mu_j)} \frac{\partial}{\partial \mu_g} \left\{ -\frac{1}{2} (x_i - \mu_g)^T \Sigma_g^{-1} (x_i - \mu_g) \right\}$$

$$= \sum_{i=1}^n \tau_i^g \Sigma_g^{-1} (x_i - \mu_g)$$

$$\begin{aligned} \tau_i^g &= p(z_i = g | x_i, \theta) \\ &= \frac{p(x_i, z_i = g | \theta)}{p(x_i | \theta)} \\ &= \frac{p(x_i | z_i = g, \theta) p(z_i = g | \theta)}{p(x_i | \theta)} \\ &= \frac{\pi_g N(x_i; \mu_g)}{\sum_{j=1}^G \pi_j N(x_i; \mu_j)} \end{aligned}$$



方程式

$$L = \sum_{i=1}^n \log \sum_{g=1}^G \pi_g N(x_i; \mu_g) + \lambda(1 - \sum_{g=1}^G \pi_g) \quad \tau_i^g = \frac{\pi_g N(x_i; \mu_g)}{\sum_{j=1}^G \pi_j N(x_i; \mu_j)}$$

$$\frac{\partial L}{\partial \pi_g} = \sum_{i=1}^n \tau_i^g / \pi_g - \lambda = 0, \quad \sum_{g=1}^G \pi_g = 1, \quad \sum_{g=1}^G \sum_{i=1}^n \tau_i^g = n \quad \text{より} \quad \pi_g = \frac{1}{n} \sum_{i=1}^n \tau_i^g$$

$$\frac{\partial L}{\partial \mu_g} = \sum_{i=1}^n \tau_i^g \Sigma_g^{-1} (x_i - \mu_g) = 0 \quad \text{より} \quad \mu_g = \frac{\sum_{i=1}^n \tau_i^g x_i}{\sum_{i=1}^n \tau_i^g}$$

非線形連立方程式だが、これは解けない。
ヒューリスティクスにより、下記のようにすればよさそうだが、果して収束するのだろうか

$$\tau_i^g \leftarrow \frac{\pi_g N(x_i; \mu_g)}{\sum_{j=1}^G \pi_j N(x_i; \mu_j)} \quad \pi_g \leftarrow \frac{1}{n} \sum_{i=1}^n \tau_i^g \quad \mu_g \leftarrow \frac{\sum_{i=1}^n \tau_i^g x_i}{\sum_{i=1}^n \tau_i^g}$$

脱線 & 参考: EMとの対応

$$\theta_{ML} = \arg \max_{\mu_g, \pi_g} LL(\mu_1, \dots, \pi_1, \dots)$$

勿論、対応するのですが、それは後講釈

$$LL(\mu_1, \dots, \pi_1, \dots) = \sum_{i=1}^n \log \sum_{g=1}^G \pi_g N(x_i; \mu_g)$$

$$\text{Lagrange関数は } L = \sum_{i=1}^n \log \sum_{g=1}^G \pi_g N(x_i; \mu_g) + \lambda(1 - \sum_{g=1}^G \pi_g)$$

$$\begin{aligned} \mu_g &\leftarrow \frac{\sum_{i=1}^n \tau_i^g x_i}{\sum_{i=1}^n \tau_i^g} & \tau_i^g &\leftarrow \frac{\pi_g N(x_i; \mu_g)}{\sum_{j=1}^G \pi_j N(x_i; \mu_j)} \\ \pi_g &\leftarrow \frac{1}{n} \sum_{i=1}^n \tau_i^g & & \end{aligned}$$

反復 (Iteration)

M step (Maximization step)

E step (Expectation step)

どのクラスに属するか？

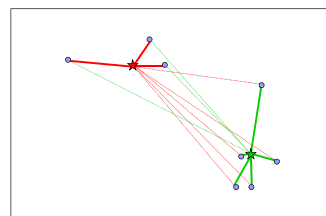
- しかし、「各 x_i は z_i 番目の分布から得られた」とする z_i は知りようがない。

- ここで、 z_i を適当に仮定する(そして、各分布のパラメータを最尤推定し、それから z_i を推定することを繰り返す)のが k-means法。

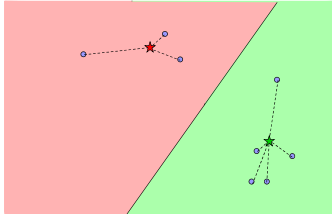
- EMアルゴリズムが考えられた

- 各点 x_i が各分布 p_g に属する確率 ($z_i = g$ となる確率)を求める
 - $\pi_g p_g(x_i)$ 正規化定数を計算する
- これらの確率をもとに、周辺尤度を最大化するパラメータを求める
 - π_g, p_g のパラメータを最尤推定する
- という反復計算。一般には、局所最適解に収束。

EM-アルゴリズムの実行イメージ

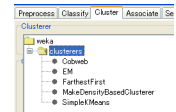


K-Means クラスタリングとの比較



EMアルゴリズムの長所と短所

- 長所
 - 統計的基盤、理論的に健全
 - 信頼性等が計算可能。検定もできる
- 短所
 - 計算が大変
 - 最近では、式さえ得られれば計算自体はかなり速くなった
 - 世の中のものに、「分布」を予め仮定するのが妥当か
 - 現実問題への適用に際し、その妥当性に疑問がある場合が多い



再: クラスタリング方法の分類

- 階層的クラスタリング:
 - ある基準に従って、階層的に対象(データ)の集合を分割していく
- 非階層的クラスタリング
 - モデル法: 各クラスに(統計的)モデルを仮定し、データに最もあうように、パラメータを決定する
 - 分割的クラスタリング: いくつかの分割を作成し、ある基準に従い評価する
 - 濃度に基づく: 濃度・密度・結合度に基づく
 - グリッド法: 多層の粒度構造に基づく

例: 文書分類はできるか?

- 実は、できない!
- クラスタリングするには、各文書をなんらかの数値ベクトルで表現しないといけない。
- どうやる?

文書クラスタリング

- 普通は次のように行う
 - 文書を単語に分ける。
 - 例: 明日は晴れでしょう。→ 明日+は+晴れ+で+しょう+。
 - どこにでもある単語(at, is, a,...)ははずす
 - なぜでしょうか?
 - 余りにも特殊な単語ははずす
 - なぜでしょうか?
 - 文書ごとに文書の特徴ベクトルを作る。各単語に番号を振り、その番号を場所とするベクトルである。各要素は、当該文書中出现する当該単語の回数(または、0/1)
 - 文書を点、この特徴ベクトルをその点の座標と考えて、クラスタリングを行う

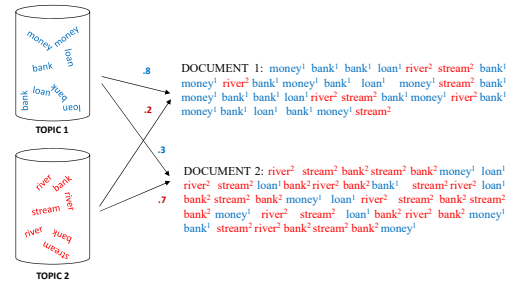
文書クラスタリング: 困ること

- ベクトルの次元が極めて高くなる
 - 短い文書(実験を行う newsgroup への投稿記事でも)かつ少ない文書数でもすぐ数千、数万になる。
 - すなわち、数千次元、数万次元。
 - 人間が図で見れるのは多くて3次元
- 実は、「高次元」というのはいろいろな問題を引き起こす要因である。
 - 計算に時間がかかる
 - (いくらメモリが 1Gとか2Gとかいっても)メモリ不足
 - 実は、ベクトルの距離の定義が難しくなる。ユークリッド距離は適さなくなってしまう

文書クラスタリング: それでもできる

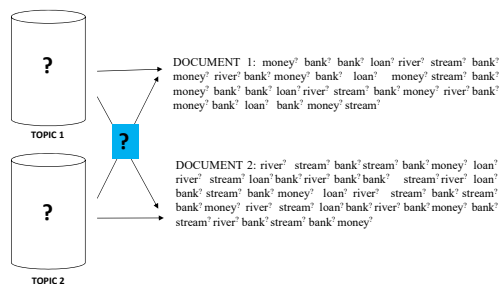
- エ夫をすることにより(さらに、コンピュータは速く、メモリが大きくなっている)、パソコンでも結構多くの文書のクラスタリングができるようになった
 - 10万文書でもできます。
- 文書を表現するベクトルの次元を落とす方法が考えられている。
 - トピック(話題)モデルという考えがある。
 - 各文書が複数個のトピックからなると考える。なお、文書中の単語は何かのトピックに属するとする。
 - 各文書を単語の(単語の出現数の)ベクトルで表すのではなく、トピックの出現頻度で表す。

脱線: トピックモデルによる、文書の生成例



http://helper.ipam.ucla.edu/publications/cog2005/cog2005_5282.ppt

脱線: トピックモデルの学習



脱線: トピックモデルの例

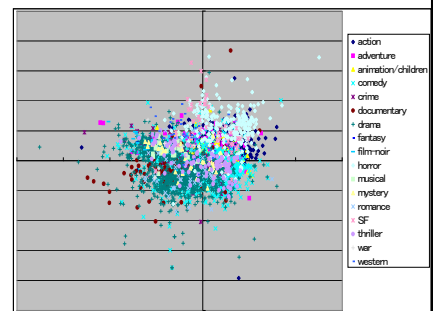
- Unigram Mixture (UM)
- pLSI/pLSA (probabilistic LSI/LSA)
- LDA (Latent Dirichlet Allocation)
 - 拡張モデルがある

次元の呪い: これは実問題の問題

- Curse of dimensionality という
- コンビニの売り上げデータを考えよう
- 一回の売り上げを一個のベクトルであらわす
 - Excel の表なら、横1行が一回の売り上げ。
- ベクトルの各要素は、各商品に対応する
 - Excel の表なら、縦1行がある商品一個に対応
- ある一年を考えるだけでも、数千・数万行が必要。
 - 次元が数千・数万のベクトル
- 数年分を管理しようと思えば、数万・数十万行
 - 次元が数万・数十万のベクトル

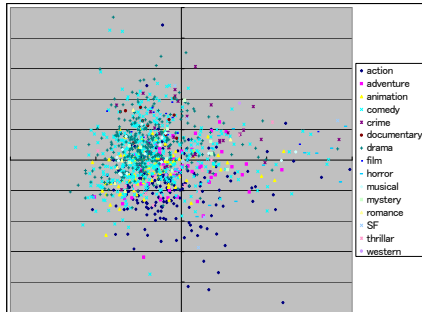
高次元の実データ例 MovieLens

- 全ユーザ
- 数量化3類適用後



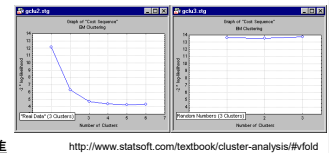
高次元の実データ例 MovieLens

- 女性、18歳
~24歳
- 数量化3類
適用後



クラスタ数の推定

- 適切なクラスタ数はどうやって推定するか？
 - cross validation (交叉検定)
 - 評価基準は、各点の対応するクラスタ中心までの距離の和
- 情報量規準を用いる
 - AIC: 赤池情報量基準
 - BIC: ベイズ情報量基準
 - Mclustではモデルの選択にBICを用いている。



Rにおける clustering その2

- パッケージとしては, stats (標準に含まれている), e1071, cluster などに含まれる
- ここでは, stats の kmeans (k-means)を試用する。
- なお、EMアルゴリズムを用いたクラスタリングには、mclustパッケージのMclustがよくつかわれる

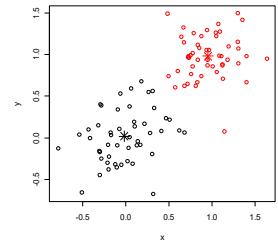
```
> # k-means in stats package
> #
> x <- rbind(matrix(rnorm(100, sd = 0.3), ncol = 2),
+             matrix(rnorm(100, mean = 1, sd = 0.3), ncol = 2))
> colnames(x) <- c("x", "y")
> (cl <- kmeans(x, 2))
K-means clustering with 2 clusters of sizes 49, 51

Cluster means:
      x      y
[1] 0.01249021 -0.0488819
[2] 0.99572877  0.8887047

Clustering vector:
[1] 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1
[32] 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1
[63] 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2
[94] 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2

Within cluster sum of squares by cluster:
[1] 8.926331 9.806389

Available components:
[1] "cluster" "centers" "withinss" "size"
> plot(x, col = cl$cluster)
> points(cl$centers, col = 1:2, pch = 8, cex=2)
>
```



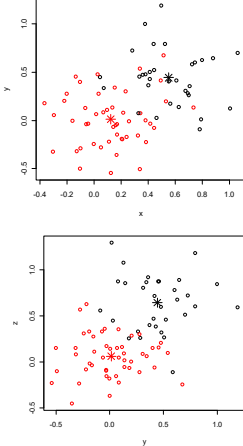
```
> x <- rbind(matrix(rnorm(120, sd = 0.3), ncol = 3),
+             matrix(rnorm(120, mean = 0.5, sd = 0.3), ncol = 3))
> colnames(x) <- c("x", "y", "z")
> (cl <- kmeans(x, 2))
K-means clustering with 2 clusters of sizes 48, 32

Cluster means:
      x      y      z
[1] 0.1210649 0.01357391 0.06235026
[2] 0.5492681 0.44262103 0.64539340

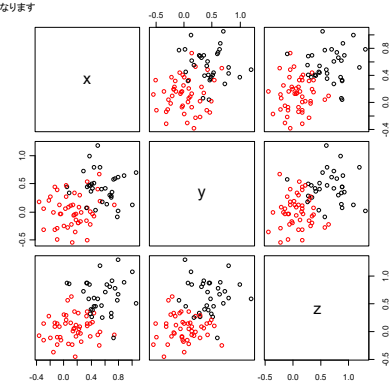
Clustering vector:
[1] 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1
[26] 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1
[51] 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2 2
[76] 2 1 2 2 2

Within cluster sum of squares by cluster:
[1] 8.754782 7.128787

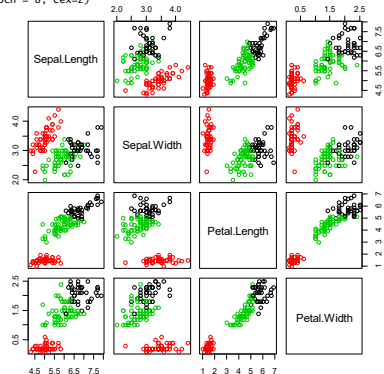
Available components:
[1] "cluster" "centers" "withinss" "size"
> plot(x, col = cl$cluster)
> points(cl$centers, col = 1:2, pch = 8, cex=2)
>
> plot(x[,2:3], col = cl$cluster)
> points(cl$centers[,2:3], col = 1:2, pch = 8, cex=2)
>
```



```
> # data.frame にすると一覧の図になります
> x <- data.frame(x)
> plot(x, col = cl$cluster)
```



```
> # iris では
> x <- iris[,1:5]
> cl <- kmeans(x, 3)
> plot(x, col = cl$cluster)
> points(cl$centers, col = 1:3, pch = 8, cex=2)
>
```



クラスタリングで分類

```
x <- iris[,1:5]
(cl <- kmeans(x, 3))
plot(x, col = cl$cluster)
points(cl$centers, col = 1:3, pch = 8, cex=2)

# クラスタリング結果をそのまま分類結果と考えたときの正解率をみてみよう
# permutations という関数を持つライブラリ
library(gtools)
# クラスタリングは(自発的な)分類と考えることができる。
# しかし、クラスタには名前がない。整数が用いられるが、真の分類とは全く別の名前である。
# そこで、クラスタと真のカテゴリとの全ての組合せを試し、正解率が一番高くなる組み合わせ
# (クラスタの並び方)を発見することを考える。

(tbl <- table(iris[,5], cl$cluster))
(res <- apply(permutations(3,3), 1, function(x) sum(diag(tbl[,x]))))
res.sorted <- sort.int(res, decreasing=T, index.return=T)
# accuracy
(acc <- res.sorted$x[1]/sum(tbl))
# そのaccuracyを与える順序(その結果のconfusion matrix)
tbl[, permutations(3,3)[res.sorted$ix[1,]]
```

R でEM

```
# EM algorithm for clustering
library(mclust)
# data sample in mlbench will be used
library(mlbench)
smly <- mlbench.smiley()
colnames(smly$x) <- c("x1", "x2")

(gm4 <- Mclust(smly$x, G=4))
mclust2dplot(smly$x, parameters=gm4$parameters,
             z=gm4$z, what="classification")
title("Clustering: 4 clusters")

dev.new()
(gm4to10 <- Mclust(smly$x, G=4:10))
mclust2dplot(smly$x, parameters=gm4to10$parameters,
             z=gm4to10$z, what="classification")
title("Clustering: best in 4 to 10 clusters")
```

K-means と EM の比較

```
# EM algorithm for clustering
library(mclust)
# data sample in mlbench will be used
library(mlbench)
smly <- mlbench.smiley()
colnames(smly$x) <- c("x1", "x2")

cl <- kmeans(smly$x, 4)
plot(smly$x, col = cl$cluster)
points(cl$centers, col = 1:4, pch = 8, cex=2)
title("Clustering: 4 clusters")

dev.new()
cl <- kmeans(smly$x, 10)
plot(smly$x, col = cl$cluster)
points(cl$centers, col = 1:10, pch = 8, cex=2)
title("Clustering: 10 clusters")
```

本日の課題1

- Animals のデータを k-means でクラスタリング分析してください。その結果を、hclust でのクラスタリング結果と比較してください。
 - Animals のデータを直接使うなら、例えば、
`x <- Animals[,2:3]` として、iris と同様にクラスタリングします。
 クラスタ数は2や3で試してください。
- 芳しくないなら、恐らく、スケールが違うからでしょうから、
`x <- myAnimals` としてみてください。
- 散布図から分かることですが、点の分布が小さい方にたくさん集まっています。それが原因かもしれません。では、対数をとってみましょう。`log(Animals[,2:3])` などとすると対数がとれます。

本日の課題2 (余裕があれば)

- `library(mlbench)` に含まれる `mlbench.spirals` を対象として、EMアルゴリズムと k-means アルゴリズムを用いたクラスタリングを行って下さい。
 - 前のスライドと同じ手続きで進めて下さい。
 - クラスタ数はいくつぐらいが妥当でしょうか？
 - それは、実際と符合しますか？
 - 符合しないとしたら、それはどうしてでしょうか？