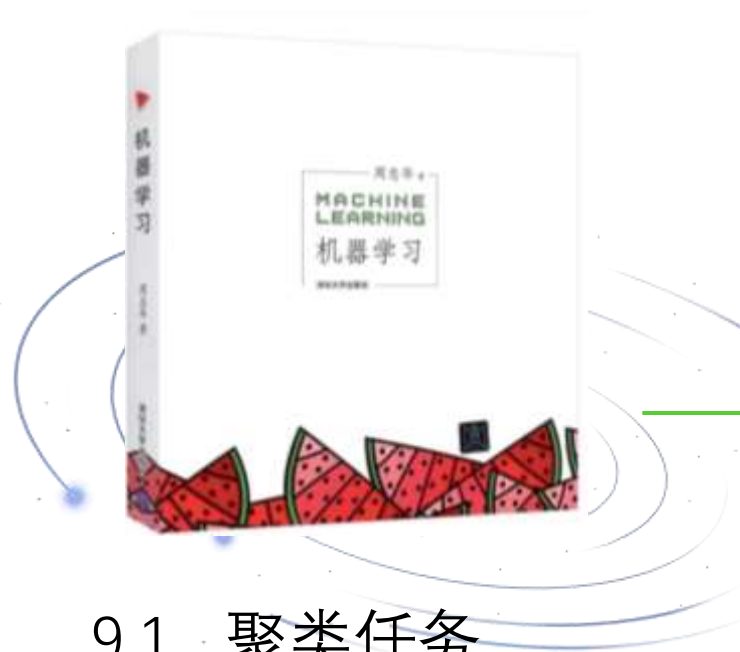




聚类

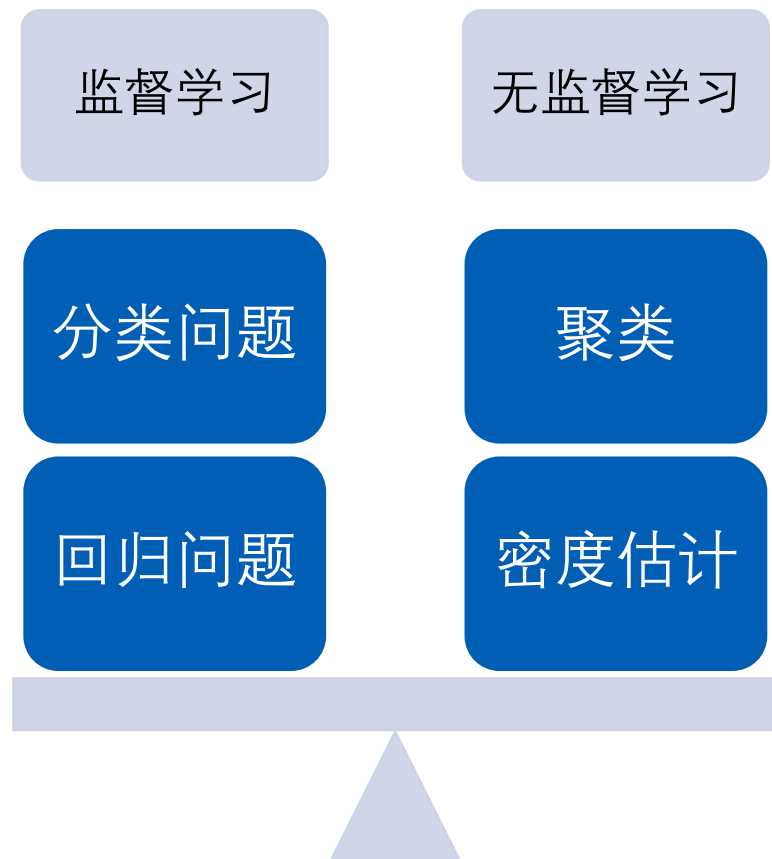
从“西瓜书”开始

9.1 聚类任务
9.3 距离计算
9.5 聚类密度

9.2 性能度量
9.4 原型聚类
9.6 层次聚类

潘贤润
2024/09/29

9.1 聚类任务



聚类试图将数据集中的样本划分为若干个通常是不相交的子集, 每个子集称为一个“簇” (cluster).

9.2 性能度量

- **validity index**

- 物以类聚
 - 簇内相似度 (intra-cluster similarity) 高
 - 簇间相似度 (inter-cluster similarity) 低

聚类性能度量大致有两类. 一类是将聚类结果与某个“参考模型” (reference model) 进行比较, 称为“外部指标” (external index); 另一类是直接考察聚类结果而不利用任何参考模型, 称为“内部指标” (internal index).

9.2.1 外部指标

对数据集 $D = \{x_1, x_2, \dots, x_m\}$, 假定通过聚类给出的簇划分为 $C = \{C_1, C_2, \dots, C_k\}$, 参考模型给出的簇划分为 $C^* = \{C_1^*, C_2^*, \dots, C_s^*\}$. 相应地, 令 λ 与 λ^* 分别表示与 C 和 C^* 对应的簇标记向量. 我们将样本两两配对考虑, 定义

4个样本
1, 2, 3, 4
 $a = |SS|$, $SS = \{(x_i, x_j) \mid \lambda_i = \lambda_j, \lambda_i^* = \lambda_j^*, i < j\}$, 同簇 (9.1)

(1,2) (1,3) (1,4)
 $b = |SD|$, $SD = \{(x_i, x_j) \mid \lambda_i = \lambda_j, \lambda_i^* \neq \lambda_j^*, i < j\}$, 假阳性 (9.2)

(2,3) (2,4) (3,4)
 $c = |DS|$, $DS = \{(x_i, x_j) \mid \lambda_i \neq \lambda_j, \lambda_i^* = \lambda_j^*, i < j\}$, 假阴性 (9.3)

4x3
2
 $d = |DD|$, $DD = \{(x_i, x_j) \mid \lambda_i \neq \lambda_j, \lambda_i^* \neq \lambda_j^*, i < j\}$, 异簇 (9.4)

其中集合 SS 包含了在 C 中隶属于相同簇且在 C^* 中也隶属于相同簇的样本对, 集合 SD 包含了在 C 中隶属于相同簇但在 C^* 中隶属于不同簇的样本对, ... 由于每个样本对 (x_i, x_j) ($i < j$) 仅能出现在一个集合中, 因此有 $a + b + c + d = m(m-1)/2$ 成立. - 共有多少对

基于式(9.1)~(9.4)可导出下面这些常用的聚类性能度量外部指标:

- Jaccard 系数(Jaccard Coefficient, 简称 JC)

集合之间相似性 $\rightarrow$ 分配到同一个簇

公式 JC = $\frac{|A \cap B|}{|A \cup B|}$ $JC = \frac{a}{a+b+c}$

- FM 指数(Fowlkes and Mallows Index, 简称 FMI)

$$FMI = \sqrt{\frac{a}{a+b} \cdot \frac{a}{a+c}}$$

- Rand 指数(Rand Index, 简称 RI)

聚类结果一致性 与 真实性
同/异簇

$$RI = \frac{2(a+d)}{m(m-1)} = \frac{a+b+c+d}{m(m-1)}$$

9.2.2 内部指标

考虑聚类结果的簇划分 $\mathcal{C} = \{C_1, C_2, \dots, C_k\}$, 定义

$$\text{avg}(C) = \frac{2}{|C|(|C| - 1)} \sum_{1 \leq i < j \leq |C|} \text{dist}(\mathbf{x}_i, \mathbf{x}_j), \quad (9.8)$$

$$\text{diam}(C) = \max_{1 \leq i < j \leq |C|} \text{dist}(\mathbf{x}_i, \mathbf{x}_j), \quad (9.9)$$

$$d_{\min}(C_i, C_j) = \min_{\mathbf{x}_i \in C_i, \mathbf{x}_j \in C_j} \text{dist}(\mathbf{x}_i, \mathbf{x}_j), \quad (9.10)$$

$$d_{\text{cen}}(C_i, C_j) = \text{dist}(\boldsymbol{\mu}_i, \boldsymbol{\mu}_j), \quad (9.11)$$

其中, $\text{dist}(\cdot, \cdot)$ 用于计算两个样本之间的距离; $\boldsymbol{\mu}$ 代表簇 C 的中心点 $\boldsymbol{\mu} = \frac{1}{|C|} \sum_{1 \leq i \leq |C|} \mathbf{x}_i$. 显然, $\text{avg}(C)$ 对应于簇 C 内样本间的平均距离, $\text{diam}(C)$ 对应于簇 C 内样本间的最远距离, $d_{\min}(C_i, C_j)$ 对应于簇 C_i 与簇 C_j 最近样本间的距离, $d_{\text{cen}}(C_i, C_j)$ 对应于簇 C_i 与簇 C_j 中心点间的距离.

基于式(9.8)~(9.11)可导出下面这些常用的聚类性能度量内部指标:

- DB 指数(Davies-Bouldin Index, 简称 DBI)

$$\text{DBI} = \frac{1}{k} \sum_{i=1}^k \max_{j \neq i} \left(\frac{\text{avg}(C_i) + \text{avg}(C_j)}{d_{\text{cen}}(\boldsymbol{\mu}_i, \boldsymbol{\mu}_j)} \right)$$

↓ 内紧
↑ 间疏

- Dunn 指数(Dunn Index, 简称 DI)

$$\text{DI} = \min_{1 \leq i \leq k} \left\{ \min_{j \neq i} \left(\frac{d_{\min}(C_i, C_j)}{\max_{1 \leq l \leq k} \text{diam}(C_l)} \right) \right\}$$

↑ 间疏
↓ 内紧

显然, DBI 的值越小越好, 而 DI 则相反, 值越大越好.

9.3 距离计算

非负性: $\text{dist}(\mathbf{x}_i, \mathbf{x}_j) \geq 0$;

自己-自己

同一性: $\text{dist}(\mathbf{x}_i, \mathbf{x}_j) = 0$ 当且仅当 $\mathbf{x}_i = \mathbf{x}_j$;

对称性: $\text{dist}(\mathbf{x}_i, \mathbf{x}_j) = \text{dist}(\mathbf{x}_j, \mathbf{x}_i)$; $i \rightleftarrows j$

直递性: $\text{dist}(\mathbf{x}_i, \mathbf{x}_j) \leq \text{dist}(\mathbf{x}_i, \mathbf{x}_k) + \text{dist}(\mathbf{x}_k, \mathbf{x}_j)$.

直接 $\leq$ 通过中介的

最常用的距离度量“闵可夫斯基距离” (Minkowski distance):

多维空间中两个点的距离 $\rightarrow$ 范数

$$\text{dist}_{\text{mk}}(\mathbf{x}_i, \mathbf{x}_j) = \left(\sum_{u=1}^n |x_{iu} - x_{ju}|^p \right)^{\frac{1}{p}}.$$

当 $p=1$ 时, 即曼哈顿距离 (Manhattan distance):

所有维度上绝对差值的总和

$$\text{dist}_{\text{man}}(\mathbf{x}_i, \mathbf{x}_j) = \|\mathbf{x}_i - \mathbf{x}_j\|_1 = \sum_{u=1}^n |x_{iu} - x_{ju}|.$$

当 $p=2$ 时, 即欧氏距离 (Euclidean distance):

两点间的直线距离

$$\text{dist}_{\text{ed}}(\mathbf{x}_i, \mathbf{x}_j) = \|\mathbf{x}_i - \mathbf{x}_j\|_2 = \sqrt{\sum_{u=1}^n |x_{iu} - x_{ju}|^2}.$$

连续属性和存在序关系的离散属性都可以直接参与计算, 因为它们都可以反映一种程度, 我们称其为“有序属性”

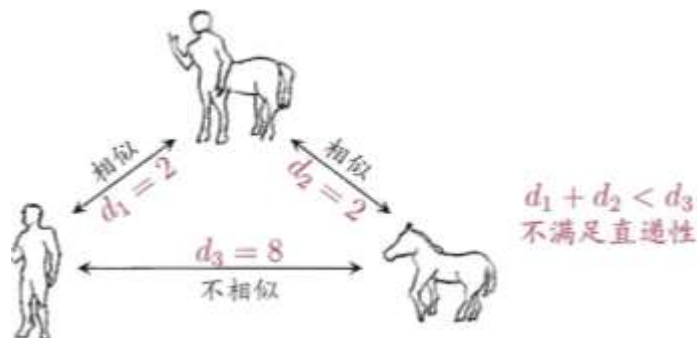
Value Difference Metric k 样本簇中为红色的数量

$$\text{VDM}_p(a, b) = \sum_{i=1}^k \left| \frac{m_{u,a,i}}{m_{u,a}} - \frac{m_{u,b,i}}{m_{u,b}} \right|^p.$$

颜色为红色的数量 为蓝色的
出现的频率差异大, VDM距离大 $\rightarrow$ 差异大

$$\text{MinkovDM}_p(\mathbf{x}_i, \mathbf{x}_j) = \left(\underbrace{\sum_{u=1}^{n_c} |x_{iu} - x_{ju}|^p}_{\text{有序属性}} + \underbrace{\sum_{u=n_c+1}^n \text{VDM}_p(x_{iu}, x_{ju})}_{\text{无序属性}} \right)^{\frac{1}{p}}. \quad (9)$$

非度量距离:



ISS
smFISH
Hamming distance

9.4 原型聚类

原型聚类也称基于**原型**的聚类，此类算法假设聚类结构能通过一组原型刻画，在现实聚类中极为常用。

- 原型是指样本空间中具有代表性的点

算法过程：

- 一般流程为先对原型进行初始化，然后对原型进行迭代更新求解。
- 采用不同的原型表示、不同的求解方式，将产生不同的算法。
- K-Means便是基于**簇中心**来实现聚类
- 混合高斯聚类则是基于**簇分布**来实现聚类

9.4.1 k均值聚类

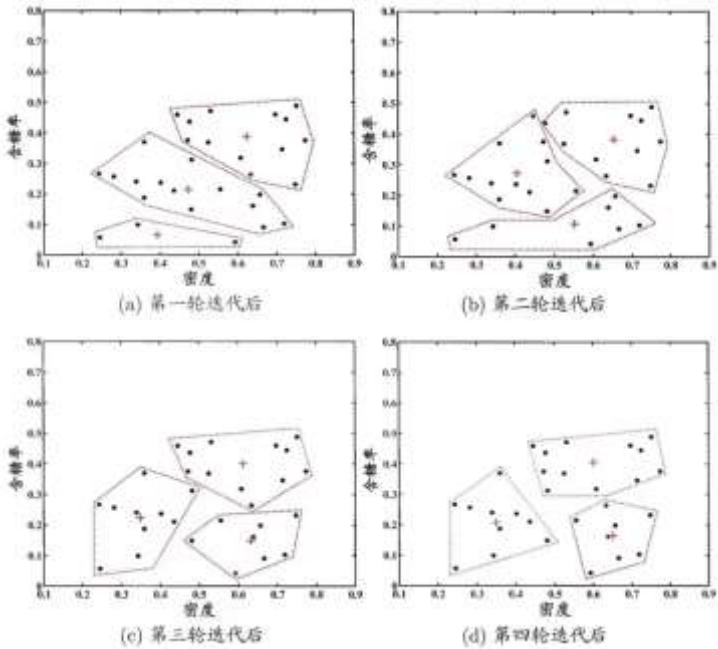
首先随机指定类中心，根据样本与类中心的远近划分类簇，接着重新计算类中心，迭代直至收敛。

给定样本集 $D = \{x_1, x_2, \dots, x_m\}$, “k 均值” (k-means) 算法针对聚类所得簇划分 $C = \{C_1, C_2, \dots, C_k\}$ 最小化平方误差

$$E = \sum_{i=1}^k \sum_{x \in C_i} \|x - \mu_i\|_2^2$$
 (9.24)

欧氏距离
簇中心点
损失误差: $\sum (y_i - \hat{y}_i)^2$
共k簇

其中 $\mu_i = \frac{1}{|C_i|} \sum_{x \in C_i} x$ 是簇 C_i 的均值向量。直观来看，式(9.24) 在一定程度上刻画了簇内样本围绕簇均值向量的紧密程度， E 值越小则簇内样本相似度越高。



输入: 样本集 $D = \{x_1, x_2, \dots, x_m\}$;
聚类簇数 k . 事先指定 k

过程:

- 1: 从 D 中随机选择 k 个样本作为初始均值向量 $\{\mu_1, \mu_2, \dots, \mu_k\}$
- 2: repeat
- 3: 令 $C_i = \emptyset$ ($1 \leq i \leq k$)
- 4: for $j = 1, 2, \dots, m$ do
- 5: 计算样本 x_j 与各均值向量 μ_i ($1 \leq i \leq k$) 的距离: $d_{ji} = \|x_j - \mu_i\|_2$;
- 6: 根据距离最近的均值向量确定 x_j 的簇标记: $\lambda_j = \arg \min_{i \in \{1, 2, \dots, k\}} d_{ji}$;
- 7: 将样本 x_j 划入相应的簇: $C_{\lambda_j} = C_{\lambda_j} \cup \{x_j\}$;
- 8: end for
- 9: for $i = 1, 2, \dots, k$ do
- 10: 计算新均值向量: $\mu'_i = \frac{1}{|C_i|} \sum_{x \in C_i} x$;
- 11: if $\mu'_i \neq \mu_i$ then
- 12: 将当前均值向量 μ_i 更新为 μ'_i
- 13: else
- 14: 保持当前均值向量不变
- 15: end if
- 16: end for
- 17: until 当前均值向量均未更新 / 或达到设置的更新次数

输出: 簇划分 $C = \{C_1, C_2, \dots, C_k\}$

欧氏距离
不含样本
m个x
表示在所有可能的类别i中，找到使得距离dji最小的那个类别i。
 $\mu_i \rightarrow \mu'_i$

9.4.2 学习向量量化

LVQ使用样本**真实类标记**辅助聚类，首先LVQ根据样本的类标记，从各类中分别随机选出一个样本作为该类簇的原型，从而组成了一个**原型特征向量组**，接着从样本集中随机挑选一个样本，计算其与原型向量组中每个向量的距离，并选取距离最小的原型向量所在的类簇作为它的划分结果，再与真实类标比较。

输入：样本集 $D = \{(x_1, y_1), (x_2, y_2), \dots, (x_m, y_m)\}$;
原型向量个数 q ，各原型向量预设的类别标记 $\{t_1, t_2, \dots, t_q\}$;
学习率 $\eta \in (0, 1)$ 。 **好坏**

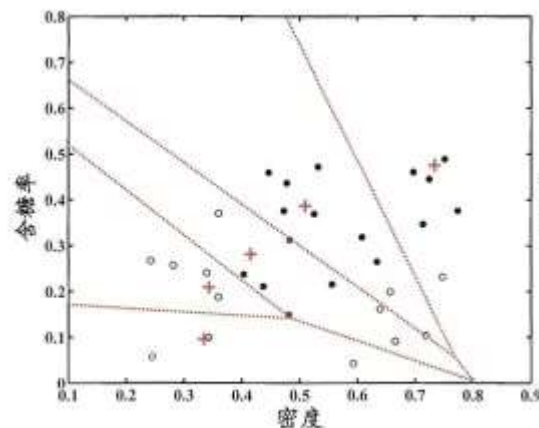
过程：

- 1: 初始化一组原型向量 $\{p_1, p_2, \dots, p_q\}$
- 2: repeat
- 3: 从样本集 D 随机选取样本 (x_j, y_j) ;
- 4: 计算样本 x_j 与 p_i ($1 \leq i \leq q$) 的距离: $d_{ji} = \|x_j - p_i\|_2$;
- 5: 找出与 x_j 距离最近的原型向量 p_{i^*} , $i^* = \arg \min_{i \in \{1, 2, \dots, q\}} d_{ji}$;
- 6: if $y_j = t_{i^*}$ then
- 7: $p' = p_{i^*} + \eta \cdot (x_j - p_{i^*})$ **让更新的 p_{i^*} 接近 x_j**
- 8: else
- 9: $p' = p_{i^*} - \eta \cdot (x_j - p_{i^*})$ $\|p' - x_j\|_2 = \|p_{i^*} + \eta \cdot (x_j - p_{i^*}) - x_j\|_2 = (1 - \eta) \cdot \|p_{i^*} - x_j\|_2$ **< 1**
- 10: end if

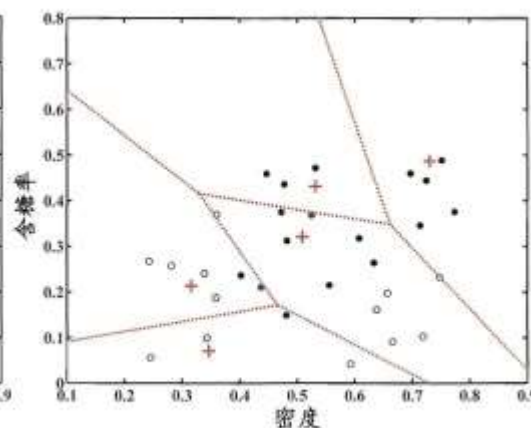
1: 将原型向量 p_{i^*} 更新为 p'

2: until 满足停止条件 **样本与 p_i 距离不大于其与 p' 的距离**

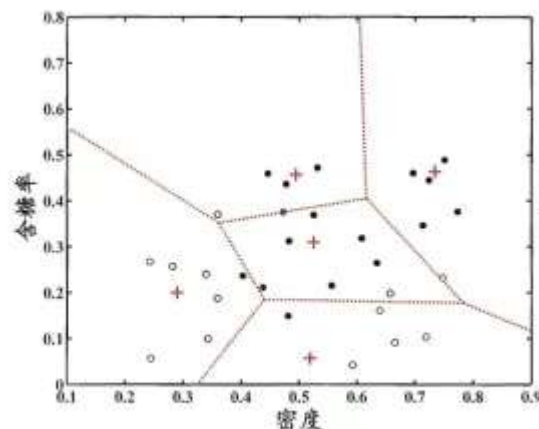
输出：原型向量 $\{p_1, p_2, \dots, p_q\}$ **中心点，而不是每个簇的样本**



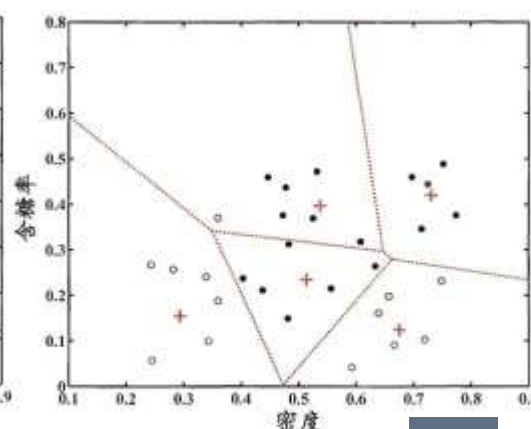
(a) 50 轮迭代后



(b) 100 轮迭代后



(c) 200 轮迭代后



(d) 400 轮迭代后

9.4.3 高斯混合聚类

对 n 维样本空间中的随机向量 x ，若 x 服从高斯分布，其概率密度函数为：

$$p(x) = \frac{1}{(2\pi)^{\frac{n}{2}} |\Sigma|^{\frac{1}{2}}} e^{-\frac{1}{2}(x-\mu)^T \Sigma^{-1}(x-\mu)}$$

其中 μ 是 n 维均值向量， Σ 是 $n \times n$ 的协方差矩阵。也可将概率密度函数记作 $p(x|\mu, \Sigma)$

$$p_{\mathcal{M}}(x) = \sum_{i=1}^k \alpha_i \cdot p(x | \mu_i, \Sigma_i),$$

α_i >0, 混合系数 $\sum_{i=1}^k \alpha_i = 1$

α_i 为选择第 i 个混合成分的概率；

对于训练集 $D = \{x_1, x_2, \dots, x_m\}$ ，现在要把 m 个样本划分为 k 个簇，即认为训练集 D 的样本是根据 k 个不同的多元高斯分布加权而得的高斯混合模型生成的。

9.4.3 高斯混合聚类

$$\begin{aligned} p_{\mathcal{M}}(z_j = i \mid \mathbf{x}_j) &= \frac{P(z_j = i) \cdot p_{\mathcal{M}}(\mathbf{x}_j \mid z_j = i)}{p_{\mathcal{M}}(\mathbf{x}_j)} \\ &= \frac{\alpha_i \cdot p(\mathbf{x}_j \mid \boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i)}{\sum_{l=1}^k \alpha_l \cdot p(\mathbf{x}_j \mid \boldsymbol{\mu}_l, \boldsymbol{\Sigma}_l)} \end{aligned}$$

全概率公式

$$\lambda_j = \arg \max_{i \in \{1, 2, \dots, k\}} \gamma_{ji} .$$

极大似然估计

$$\begin{aligned} LL(D) &= \ln \left(\prod_{j=1}^m p_M(x_j) \right) \\ &= \sum_{j=1}^m \ln \left(\sum_{i=1}^k \alpha_i \cdot p(x_j \mid \mu_i, \Sigma_i) \right) \end{aligned}$$

$$\text{令 } \frac{\partial LL(D)}{\partial \mu_i} = 0 \text{ 得 } \mu_i = \frac{\sum_{j=1}^m \gamma_{ji} x_j}{\sum_{j=1}^m \gamma_{ji}}$$

$$\text{令 } \frac{\partial LL(D)}{\partial \Sigma_i} = 0 \text{ 得 } \Sigma_i = \frac{\sum_{j=1}^m \gamma_{ji} (x_j - \mu_i)(x_j - \mu_i)^T}{\sum_{j=1}^m \gamma_{ji}}$$

拉格朗日乘子

$$\alpha_i = \frac{1}{m} \sum_{j=1}^m \gamma_{ji}$$

9.4.3 高斯混合聚类

输入: 样本集 $D = \{x_1, x_2, \dots, x_m\}$; m ↑

高斯混合成分个数 k .

EM 算法:
(期望最大算法)

过程:

1: 初始化高斯混合分布的模型参数 $\{(\alpha_i, \mu_i, \Sigma_i) \mid 1 \leq i \leq k\}$

2: repeat

3: for $j = 1, 2, \dots, m$ do

4: 根据式(9.30)计算 x_j 由各混合成分生成的后验概率, 即

$\gamma_{ji} = p_{\mathcal{M}}(z_j = i \mid x_j) \quad (1 \leq i \leq k)$ E步 → 估计参数 (期望步)

5: end for

6: for $i = 1, 2, \dots, k$ do

7: 计算新均值向量: $\mu'_i = \frac{\sum_{j=1}^m \gamma_{ji} x_j}{\sum_{j=1}^m \gamma_{ji}}$;

8: 计算新协方差矩阵: $\Sigma'_i = \frac{\sum_{j=1}^m \gamma_{ji} (x_j - \mu'_i)(x_j - \mu'_i)^T}{\sum_{j=1}^m \gamma_{ji}}$;

9: 计算新混合系数: $\alpha'_i = \frac{\sum_{j=1}^m \gamma_{ji}}{m}$;

(极大步)

10: end for

11: 将模型参数 $\{(\alpha_i, \mu_i, \Sigma_i) \mid 1 \leq i \leq k\}$ 更新为 $\{(\alpha'_i, \mu'_i, \Sigma'_i) \mid 1 \leq i \leq k\}$

M步 → 调整参数

12: until 满足停止条件

13: $C_i = \emptyset \quad (1 \leq i \leq k)$ 每个簇集合为空

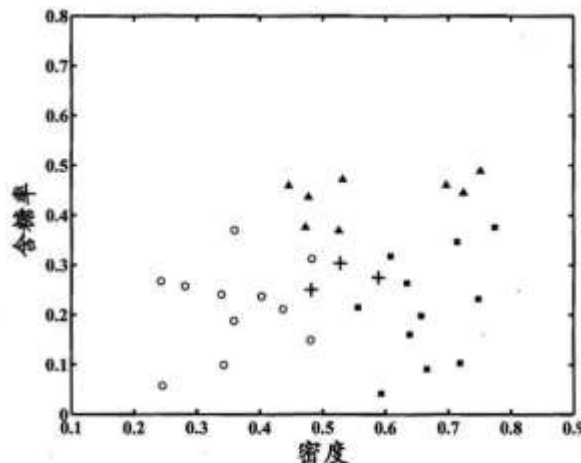
14: for $j = 1, 2, \dots, m$ do

15: 根据式(9.31)确定 x_j 的簇标记 λ_j ;

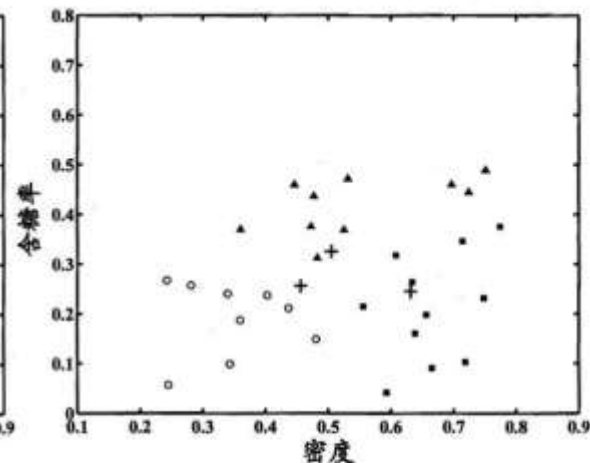
16: 将 x_j 划入相应的簇: $C_{\lambda_j} = C_{\lambda_j} \cup \{x_j\}$

17: end for

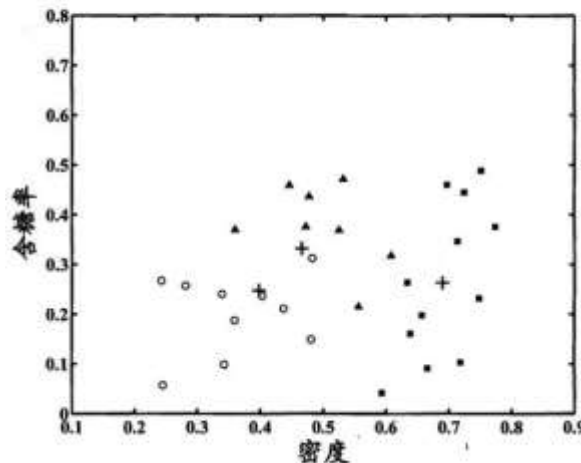
输出: 簇划分 $\mathcal{C} = \{C_1, C_2, \dots, C_k\}$



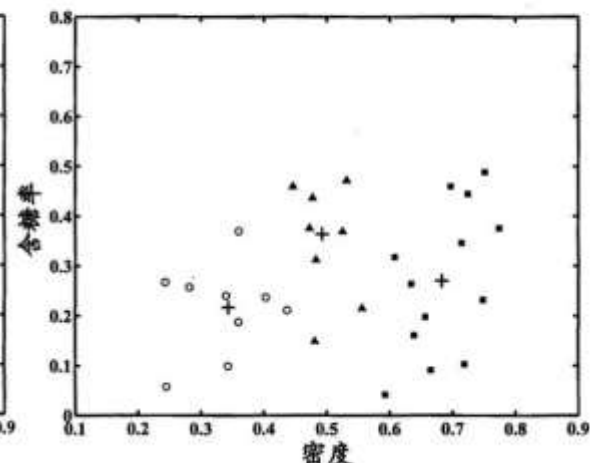
(a) 5 轮迭代后



(b) 10 轮迭代后



(c) 20 轮迭代后



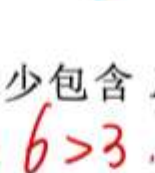
(d) 50 轮迭代后

9.5 密度聚类

密度聚类也称为“基于密度的聚类”(density-based clustering)。此类算法假设聚类结构能通过样本分布的紧密程度来确定。通常情况下，密度聚类算法从样本密度的角度来考察样本之间的可连接性，并基于可连接样本不断扩展聚类簇来获得最终的聚类结果。

DBSCAN算法：

ϵ -邻域：对 $x_j \in D$ ，其 ϵ -邻域包含样本集 D 中与 x_j 的距离不大于 ϵ 的样本，即 $N_\epsilon(x_j) = \{x_i \in D \mid \text{dist}(x_i, x_j) \leq \epsilon\}$ ；
 圈内的，是 x_j 的邻域

核心对象(core object)：若 x_j 的 ϵ -邻域至少包含 MinPts 个样本，即 $|N_\epsilon(x_j)| \geq \text{MinPts}$ ，则 x_j 是一个核心对象；
 $6 > 3$ ， x_j 是一个核心对象

密度直达(directly density-reachable)：若 x_j 位于 x_i 的 ϵ -邻域中，且 x_i 是核心对象，则称 x_j 由 x_i 密度直达；
 x_j 由 x_i 密度直达

密度可达(density-reachable)：对 x_i 与 x_j ，若存在样本序列 $p_1, p_2, \dots, p_n$ ，其中 $p_1 = x_i$ ， $p_n = x_j$ 且 p_{i+1} 由 p_i 密度直达，则称 x_j 由 x_i 密度可达；
 通过多次直达即可达

密度相连(density-connected)：对 x_i 与 x_j ，若存在 x_k 使得 x_i 与 x_j 均由 x_k 密度可达，则称 x_i 与 x_j 密度相连。

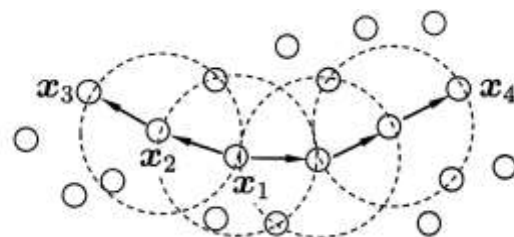


图 9.8 DBSCAN 定义的基本概念($\text{MinPts} = 3$): 虚线显示出 ϵ -邻域, x_1 是核心对象, x_2 由 x_1 密度直达, x_3 由 x_1 密度可达, x_3 与 x_4 密度相连。

简单来理解DBSCAN是：找出一个核心对象所有密度可达的样本集合形成簇

9.5 密度聚类

输入: 样本集 $D = \{x_1, x_2, \dots, x_m\}$;
邻域参数 $(\epsilon, MinPts)$.

过程:

1: 初始化核心对象集合: $\Omega = \emptyset$

2: for $j = 1, 2, \dots, m$ do

3: 确定样本 x_j 的 ϵ -邻域 $N_\epsilon(x_j)$; *找出所有*

4: if $|N_\epsilon(x_j)| \geq MinPts$ then *核心对象*

5: 将样本 x_j 加入核心对象集合: $\Omega = \Omega \cup \{x_j\}$

6: end if

7: end for

8: 初始化聚类簇数: $k = 0$ *不用指定*

9: 初始化未访问样本集合: $\Gamma = D$

10: while $\Omega \neq \emptyset$ do *逐一访问所有样本*

11: 记录当前未访问样本集合: $\Gamma_{old} = \Gamma$;

12: 随机选取一个核心对象 $o \in \Omega$, 初始化队列 $Q = \langle o \rangle$;

13: $\Gamma = \Gamma \setminus \{o\}$; *减0 直到Γ为空*

14: while $Q \neq \emptyset$ do

15: 取出队列 Q 中的首个样本 q ;

16: if $|N_\epsilon(q)| \geq MinPts$ then

17: 令 $\Delta = N_\epsilon(q) \cap \Gamma$; *删除已经被访问到的0 找到密度直达*

18: 将 Δ 中的样本加入队列 Q ;

19: $\Gamma = \Gamma \setminus \Delta$; *减去已被访问到的0*

20: end if

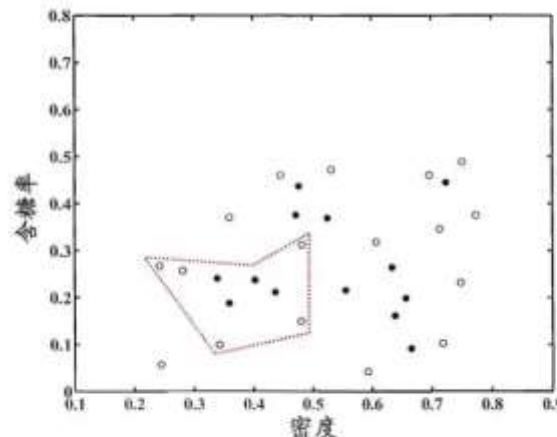
21: end while

22: $k = k + 1$, 生成聚类簇 $C_k = \Gamma_{old} \setminus \Gamma$; *核心对象o的*

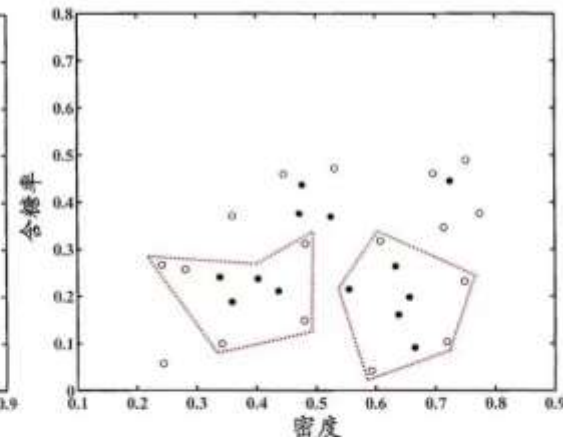
23: $\Omega = \Omega \setminus C_k$ *已访问全 未访问 密度直达 样本集合*

24: end while

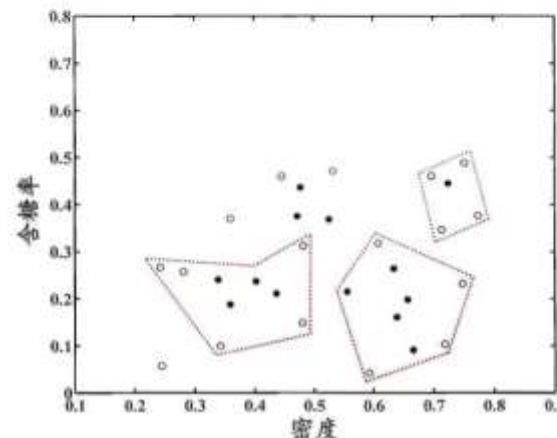
输出: 簇划分 $C = \{C_1, C_2, \dots, C_k\}$



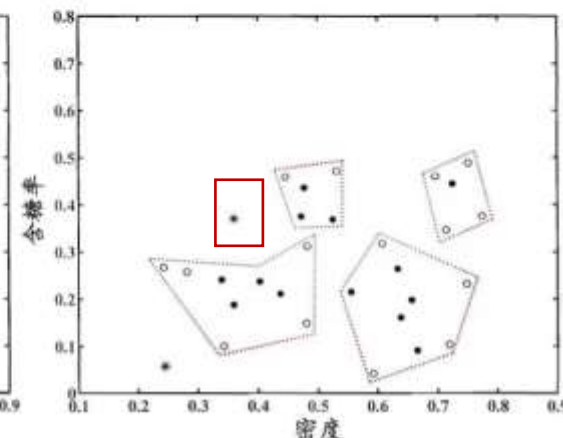
(a) 生成聚类簇 C_1



(b) 生成聚类簇 C_2



(c) 生成聚类簇 C_3



(d) 生成聚类簇 C_4

9.6 层次聚类

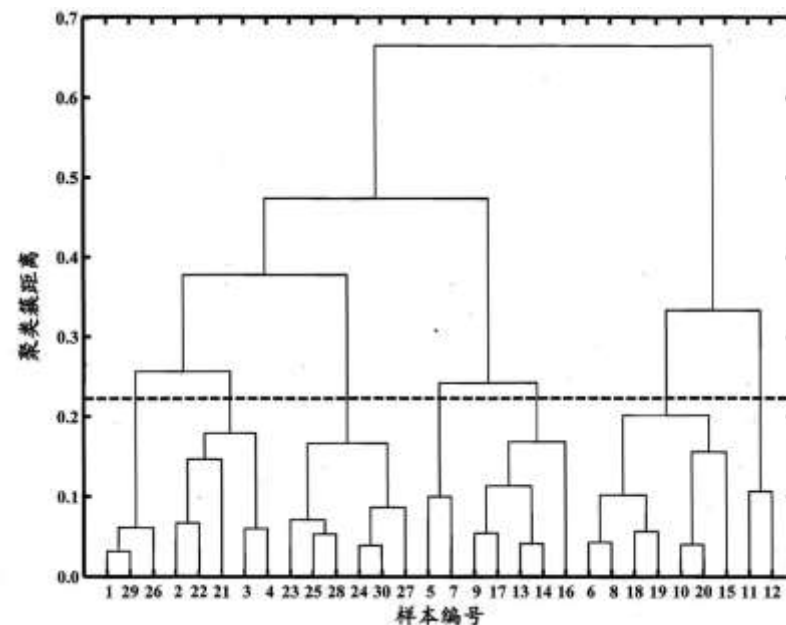
层次聚类在不同层次对数据集进行划分，从而形成树形的聚类结构。数据集划分既可采用“自底向上”的聚合策略，也可采用“自顶向下”的分拆策略。

AGNES算法（自底向上的层次聚类算法）：首先，将样本中的每一个样本看做一个初始聚类簇，然后在算法运行的每一步中找出距离最近的两个聚类簇进行合并，该过程不断重复，直到达到预设的聚类簇的个数。

最小距离: $d_{\min}(C_i, C_j) = \min_{x \in C_i, z \in C_j} \text{dist}(x, z)$, 最近样本

最大距离: $d_{\max}(C_i, C_j) = \max_{x \in C_i, z \in C_j} \text{dist}(x, z)$, 最远样本

平均距离: $d_{\text{avg}}(C_i, C_j) = \frac{1}{|C_i||C_j|} \sum_{x \in C_i} \sum_{z \in C_j} \text{dist}(x, z)$ 所有样本



9.6 层次聚类

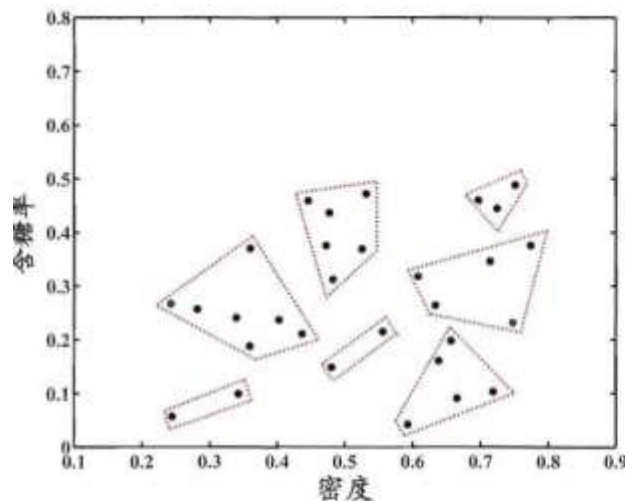
输入: 样本集 $D = \{x_1, x_2, \dots, x_m\}$;

聚类簇距离度量函数 d ;

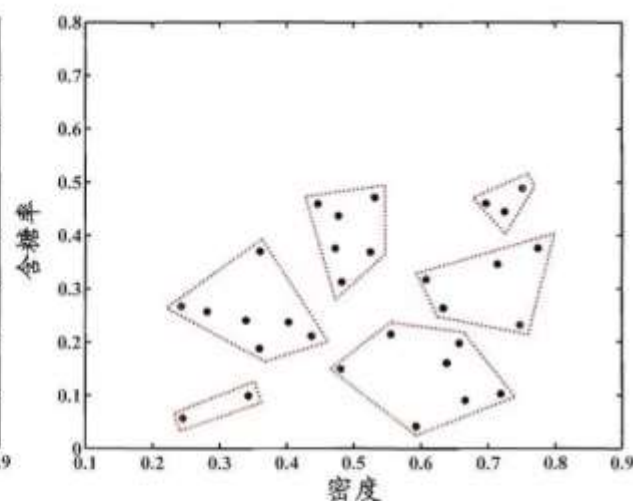
聚类簇数 k .

过程:

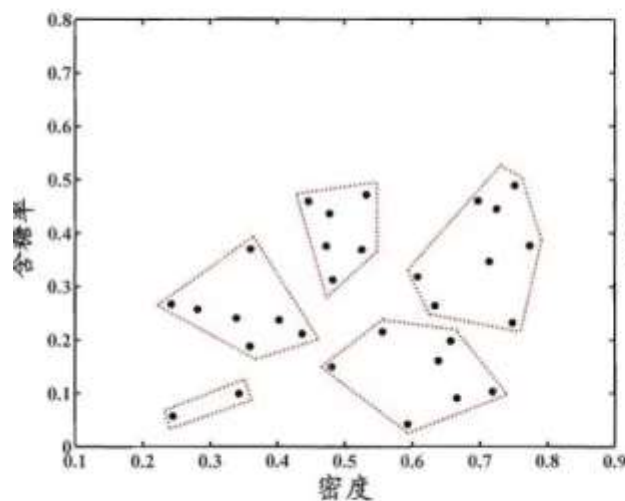
```
1: for  $j = 1, 2, \dots, m$  do
2:    $C_j = \{x_j\}$  m个簇
3: end for
4: for  $i = 1, 2, \dots, m$  do
5:   for  $j = 1, 2, \dots, m$  do
6:      $M(i, j) = d(C_i, C_j)$ ; 两两样本间距离
7:      $M(j, i) = M(i, j)$ 
8:   end for
9: end for
10: 设置当前聚类簇个数:  $q = m$ 
11: while  $q > k$  do
12:   找出距离最近的两个聚类簇  $C_{i^*}$  和  $C_{j^*}$ ;
13:   合并  $C_{i^*}$  和  $C_{j^*}$ :  $C_{i^*} = C_{i^*} \cup C_{j^*}$ ;
14:   for  $j = j^* + 1, j^* + 2, \dots, q$  do
15:     将聚类簇  $C_j$  重编号为  $C_{j-1}$ 
16:   end for
17:   删除距离矩阵  $M$  的第  $j^*$  行与第  $j^*$  列;
18:   for  $j = 1, 2, \dots, q - 1$  do
19:      $M(i^*, j) = d(C_{i^*}, C_j)$ ;
20:      $M(j, i^*) = M(i^*, j)$ 
21:   end for
22:    $q = q - 1$ 
23: end while
输出: 簇划分  $\mathcal{C} = \{C_1, C_2, \dots, C_k\}$ 
```



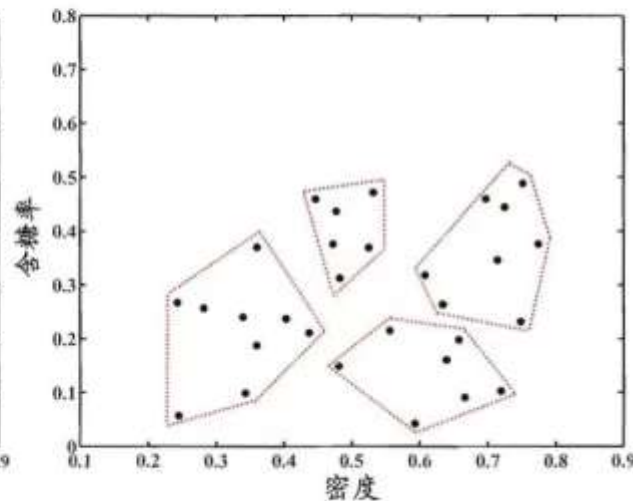
(a) 聚类簇数 $k = 7$



(b) 聚类簇数 $k = 6$



(c) 聚类簇数 $k = 5$



(d) 聚类簇数 $k = 4$

感谢倾听

See you next.

