

一种基于邻域协同表达的分类方法

徐苏平¹ 杨习贝¹ 于化龙¹ 於东军²

(江苏科技大学计算机科学与工程学院 镇江 212003)¹

(南京理工大学计算机科学与工程学院 南京 210094)²

摘要 邻域粗糙集模型中,随着信息粒尺寸的增长,基于多数投票原则的邻域分类器(NC)容易对未知样本的类别产生误判。为了缓解该问题,在协同表达分类(CRC)思想的基础上,提出了一种基于邻域协同表达的分类方法,即邻域协同分类器(NCC)。NCC首先借助邻域粗糙集模型对分类学习任务进行特征选择,然后找出被选特征下未知样本的邻域空间,最后在邻域空间内采用协同表达来代替多数投票原则,找出与未知样本具有最小重构误差的类别作为预测的类别标记。在4组UCI数据集上的实验结果表明:1)与NC相比,所提NCC在大尺寸信息粒下获得了较为满意的分类效果;2)与CRC相比,所提NCC在保持良好分类精度的同时,极大地降低了字典样本的规模,进而提高了分类的效率。

关键词 分类,协同表达,特征选择,邻域,粗糙集

中图法分类号 TP18 文献标识码 A DOI 10.11896/j.issn.1002-137X.2017.09.044

Neighborhood Collaborative Representation Based Classification Method

XU Su-ping¹ YANG Xi-bei¹ YU Hua-long¹ YU Dong-jun²

(School of Computer Science and Engineering, Jiangsu University of Science and Technology, Zhenjiang 212003, China)¹

(School of Computer Science and Engineering, Nanjing University of Science and Technology, Nanjing 210094, China)²

Abstract In the neighborhood rough set model, with the increasing of the size of information granules, the majority voting rule based neighborhood classifier (NC) is easy to misjudge the classes of unknown samples. To remedy this deficiency, based on the idea of collaborative representation based classification (CRC), we proposed a neighborhood collaborative representation based classification method, namely, the neighborhood collaborative classifier (NCC). NCC firstly performs feature selection in the classification learning task with neighborhood rough set model, and then finds the neighborhood space of unknown sample under selected features. Finally, instead of the majority voting rule in the neighborhood space, NCC judges the class of unknown sample with the collaborative representation, which considers the class with the minimal reconstruction error for unknown sample as the predicted category. Experimental results on 4 UCI data sets show that compared with NC, the proposed NCC achieves satisfactory performance in larger information granules and compared with CRC, and the proposed NCC greatly reduces the size of the dictionary while maintaining good classification accuracy, and improves the efficiency of classification.

Keywords Classification, Collaborative representation, Feature selection, Neighborhood, Rough set

1 引言

波兰学者 Pawlak 于 1982 年提出的粗糙集理论^[1-2]是一种处理不完备性、不确定性问题的强有力的数学工具。近些年,粗糙集理论已被广泛应用于数据挖掘^[3]、决策支持^[4]、特征选择^[5-6]等众多研究领域。Pawlak 粗糙集模型中最基本的概念是等价关系,它可直接由一系列的名义型特征诱导产生。

然而,在真实世界中存在着大量的数值型特征,为了让经典的粗糙集方法更好地适应现实应用, Hu 等人^[7-9]借助 3 种不同的距离函数提出了邻域粗糙集模型。在 Hu 的邻域粗糙集模型中,邻域信息粒可由名义型特征和数值型特征共同诱导产生。进一步地, Hu 以邻域粗糙集模型为基本框架实现了一种基于多数原则的邻域分类器(Neighborhood Classifiers, NC)。值得注意的是,随着邻域信息粒尺寸的增长,大量低相

到稿日期:2016-08-08 返修日期:2016-11-26 本文受国家自然科学基金(61572242, 61503160, 61305058, 61373062), 中国博士后科学基金(2014M550293), 江苏省普通高校学术学位研究生科研创新计划项目(KYLX16_0505)资助。

徐苏平(1991—),男,硕士生,主要研究方向为粗糙集理论、机器学习, E-mail: supingxu@yahoo.com; 杨习贝(1980—),男,博士后,副教授,硕士生导师,主要研究方向为粗糙集理论、粒计算与机器学习, E-mail: zhenjiangyangxibei@163.com; 于化龙(1982—),男,博士,副教授,硕士生导师,主要研究方向为机器学习、生物信息学, E-mail: yuhualong@just.edu.cn; 於东军(1975—),男,博士,教授,博士生导师,主要研究方向为生物信息学, E-mail: njyudj@njjust.edu.cn。

似度的样本落入未知样本的邻域空间中,这将导致基于多数投票原则的邻域决策容易对未知样本的类别产生误判。为此,NC的分类准确率会随着邻域决策错误率的增大而逐渐下降。

最近,一些基于稀疏性的模式分类方法引起了研究人员的广泛关注。其中最具代表性的是Wright等人^[10]提出的基于稀疏表达的分类方法(Sparse Representation Based Classification, SRC)。在SRC的基础上,为了降低稀疏求解过程的计算复杂性,Zhang等人^[11]通过分析SRC的工作机理证明了协同表达对SRC的良好分类效果起到了决定性的作用,并提出了基于正则化最小二乘解的协同表达分类方法(Collaborative Representation Based Classification, CRC)。SRC与CRC均做如下假设:任一未知样本能被视作与该未知样本同类的训练样本的线性组合。本文基于上述假设提出一种邻域协同分类器(Neighborhood Collaborative Classifier, NCC)。NCC首先借助邻域粗糙集模型对分类学习任务进行特征选择,然后找出被选特征下未知样本的邻域空间,最后在邻域空间内采用协同表达的思想对未知样本的类别进行预测。在4组UCI数据集上的实验结果表明:与NC相比,本文提出的NCC在大尺寸信息粒下获得了良好的分类效果;此外,较CRC而言,NCC以局部的方式极大地降低了字典样本的规模,进而提高了分类的效率。

2 相关工作

2.1 邻域粗糙集方法

形式化地,分类学习任务可视为四元组 $\langle U, A, V, f \rangle$,其中, $U = \{x_1, x_2, \dots, x_n\}$ 是由 n 个样本构成的非空有限集合,称为论域; $A = \{a_1, a_2, \dots, a_m\}$ 用以描述论域中所有样本的特征集合; V_a 是特征 a 的值域; $f: U \times A \rightarrow V$ 为信息函数。特别地,若 $A = C \cup D$, C 为条件特征集合, D 为决策特征,则分类学习任务 $\langle U, A, V, f \rangle$ 也可被称为决策系统(Decision System, DS)。

定义1 给定一论域 $U, \forall x_i, x_j, x_k \in U, B \subseteq C$,特征集合 B 下不同样本间的度量函数 Δ_B 定义如下:

- (1) $\Delta_B(x_i, x_j) \geq 0, \Delta_B(x_i, x_j) = 0$ 当且仅当 $x_i = x_j$;
- (2) $\Delta_B(x_i, x_j) = \Delta_B(x_j, x_i)$;
- (3) $\Delta_B(x_i, x_j) \leq \Delta_B(x_i, x_k) + \Delta_B(x_k, x_j)$ 。

显然, Δ_B 满足非负性、对称性和三角不等式。目前已有大量的度量函数被广泛使用^[12],如曼哈顿距离、欧氏距离、高斯距离等。

定义2 给定任意样本 $x_i \in U, B \subseteq C$,特征集合 B 下 x_i 的邻域 $\delta_B(x_i)$ 如下定义:

$$\delta_B(x_i) = \{x_j \mid x_j \in U, \Delta_B(x_i, x_j) \leq \delta\} \quad (1)$$

其中, δ 是邻域半径。

$\delta_B(x_i)$ 是由样本 x_i 在特征集合 B 下诱导的邻域信息粒,其大小依赖于邻域半径 δ 的选取, δ 越大,落入到 x_i 的邻域空间中的样本越多。一系列的邻域信息粒 $\{\delta_B(x_i) \mid x_i \in U\}$ 构成了基本概念的集合,其覆盖了整个论域。若特征集合 B 诱导了一系列的邻域信息粒,则称决策系统为邻域决策系统(Neighborhood Decision System, NDS)。

定义3 给定一邻域决策系统 $NDS = \langle U, C \cup D, V, f \rangle$, $\delta_B(x_i)$ 是由样本 x_i 在特征集合 $B \subseteq C$ 下诱导的邻域信息粒, $\forall S \subseteq U$,特征集合 B 下 S 的下、上近似集合分别有如下定义:

$$\underline{R}_B(S) = \{x_i \mid \delta_B(x_i) \subseteq S, x_i \in U\} \quad (2)$$

$$\overline{R}_B(S) = \{x_i \mid \delta_B(x_i) \cap S \neq \emptyset, x_i \in U\} \quad (3)$$

$[\underline{R}_B(S), \overline{R}_B(S)]$ 即为 S 的邻域粗糙集。

定义4 给定一邻域决策系统 $NDS = \langle U, C \cup D, V, f \rangle$, $U/IND(D) = \{\pi d_1, \pi d_2, \dots, \pi d_p\}$ 是由决策特征 D 诱导的论域上的划分, $\forall B \subseteq C$,决策特征 D 相对于条件特征集合 B 的邻域依赖度可用 $\gamma(B, D)$ 表示,其定义如下:

$$\gamma(B, D) = \frac{|\bigcup_{i=1}^p \underline{R}_B(\pi d_i)|}{|U|} \quad (4)$$

其中, $|\cdot|$ 表示集合的基数, $\forall \pi d_i \in U/IND(D), \pi d_i$ 为决策类。

$\gamma(B, D)$ 反映了由特征集合 B 诱导的粒化空间对决策 D 上各决策类的近似能力。若 $\gamma(B, D) = 1$,则 D 完全依赖于 B ,且 NDS 是一致的;否则, D 是 γ -依赖于 B 且 NDS 是不一致的。文献[7-8]证明了邻域决策系统中的邻域依赖度具有单调性,即若 $B_1 \subseteq B_2, \gamma(B_1, D) \leq \gamma(B_2, D)$ 成立,则此性质有助于特征选择算法^[5-6, 13-14]中贪心搜索策略的实施。

定义5 给定一邻域决策系统 $NDS = \langle U, C \cup D, V, f \rangle$, $x_i \in U, \delta_B(x_i)$ 是由样本 x_i 在特征集合 $B \subseteq C$ 下诱导的邻域信息粒, $\{P(d_1 | \delta_B(x_i)), P(d_2 | \delta_B(x_i)), \dots, P(d_p | \delta_B(x_i))\}$ 是样本 x_i 邻域内的样本归属于每一种决策类别 d_j 的概率集合,则样本 x_i 的邻域决策(Neighborhood Decision, ND)定义如下:

$$ND(x_i) = d_k, \text{ 当且仅当 } P(d_k | \delta_B(x_i)) = \max_j P(d_j | \delta_B(x_i)) \quad (5)$$

其中, $P(d_j | \delta_B(x_i)) = n_j / N, j = 1, 2, \dots, p$ 。 N 是 $\delta_B(x_i)$ 中样本的总数, n_j 是 $\delta_B(x_i)$ 中决策特征为 d_j 的样本的数量。

$ND(x_i)$ 是指在 x_i 的邻域内根据多数投票原则安排给样本 x_i 的决策类别,图1形象地展示了邻域决策的过程。

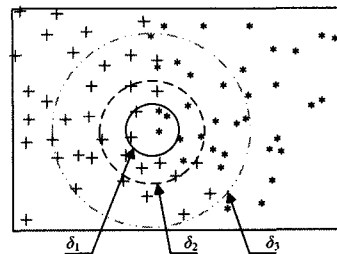


图1 邻域决策示意图

值得注意的是,随着邻域半径 δ 的增加,样本 x_i 的邻域空间内低相似度的样本也将增多,从而导致邻域决策的性能下降。图1中 $\delta_1 < \delta_2 < \delta_3$,可以发现,当半径为 δ_1 时,样本 x_1 的邻域内存在3个同类样本和2个异类样本,根据多数投票原则,样本 x_1 将正确分类;然而,随着半径增大至 δ_2 时,样本 x_1 的邻域内存在6个同类样本和9个异类样本,此时样本 x_1 将错误分类;当邻域半径继续增加时,样本 x_1 的邻域空间内的异类样本将越来越多,从而直接干扰邻域分类器(NC)对未知样本的类别判断。邻域决策方法的有效性可用邻域决策错

误率(Neighborhood Decision Error Rate, NDER)进行评估。

定义 6 邻域决策错误率的定义如下:

$$NDER = \frac{\sum_{i=1}^n \phi(d(x_i) | ND(x_i))}{n} \quad (6)$$

并且

$$\phi(d(x_i) | ND(x_i)) = \begin{cases} 0, & d(x_i) = ND(x_i) \\ 1, & d(x_i) \neq ND(x_i) \end{cases} \quad (7)$$

其中, n 是样本的数量, $d(x_i)$ 是 x_i 的真实的决策类别。

NDER 反映了在相应的特征空间下, 采用邻域决策规则时样本被错误分类的比例。

2.2 协同表达分类器

假若 $X = [X_1, X_2, \dots, X_c] \in \mathbb{R}^{n \times n}$ 是具有 c 种不同类别的 n 个样本构成的训练集合, 其中 $X_i = [x_{i,1}, x_{i,2}, \dots, x_{i,n_i}]$ 代表 n_i 个第 i 类 ($i \leq c$) 的训练样本。任一未知样本 $y \in \mathbb{R}^n$ 都能按照 $y \approx X \cdot \alpha$ 进行稀疏编码, 其中 $\alpha = [\alpha_1, \alpha_2, \dots, \alpha_c]$, $\alpha_i \in \mathbb{R}^{n_i}$ 为第 i 类的编码向量。为了使 X 重构 y 的过程具有较低的计算复杂度, Zhang 等人^[11] 提出了基于正则化最小二乘法的协同表达方式, 其形式如下:

$$(\hat{\alpha}) = \arg \min_{\alpha} \{ \|y - X \cdot \alpha\|_2^2 + \lambda \|\alpha\|_2^2 \} \quad (8)$$

其中, λ 为正则化参数。由式(8)不难求得:

$$(\hat{\alpha}) = (X^T X + \lambda \cdot I)^{-1} X^T y \quad (9)$$

其中, I 为单位矩阵。通过系数 $\hat{\alpha}$ 诱导的分类与每种类别的表示残差 $\|y - X_i \cdot \hat{\alpha}_i\|_2$ 紧密关联, 由于 $\|\hat{\alpha}_i\|_2$ 也具有一些对分类有作用的鉴别信息, 文献[11]将两者进行了融合。形式化地, 协同表达分类器设计如下。

算法 1 协同表达分类器(CRC)

输入: 训练集 X , 未知样本 y , 正则化参数 λ , 单位矩阵 I

输出: 预测的 y 的类别

1. 对 X 的每一列采用 l_2 -范式进行归一化;
2. 用 X 对 y 进行编码:

$$(\hat{\alpha}) = (X^T X + \lambda \cdot I)^{-1} X^T y;$$

3. 对每类样本计算正则化残差:

$$r_i(y) = \frac{\|y - X_i \cdot \hat{\alpha}_i\|_2}{\|\hat{\alpha}_i\|_2};$$

4. 将最小 $r_i(y)$ 所在的类别 i 作为 y 的类别。

CRC 的基本流程是: 1) 在训练集上计算未知样本的协同表达系数的解; 2) 对每类训练样本分别计算重构误差; 3) 选择具有最小重构误差的类别作为未知样本的类别标记。

3 分类学习算法

3.1 特征选择方法

分类学习任务中常常会存在大量的冗余信息, 这不仅会增加分类过程的时间开销, 而且会降低分类学习算法的性能。为了缓解上述问题, 一种有效的方法是特征选择。特征选择旨在找出与原始条件特征集合具有相同依赖度的最小特征子集。下面将讨论邻域粗糙集框架下的特征选择方法。

定义 7 给定一邻域决策系统 $NDS = \langle U, C \cup D, V, f \rangle$,

$\forall B \subseteq C, B$ 为 C 的一个约简, 当且仅当:

- (1) $\gamma(B, D) = \gamma(C, D)$;
- (2) $\forall B' \subset B, \gamma(B', D) \neq \gamma(B, D)$ 。

由定义 7 可知, 条件特征 C 的约简是 C 的保持邻域依赖度不变的最小子集。

定义 8 给定一邻域决策系统 $NDS = \langle U, C \cup D, V, f \rangle$, $\forall B \subseteq C, \forall a_i \in C - B$, 特征 a_i 的重要度定义如下:

$$Sig(a_i, B, D) = \gamma(B \cup a_i, D) - \gamma(B, D) \quad (10)$$

$Sig(a_i, B, D)$ 反映了当 a_i 被加入到 B 时邻域依赖度的变化程度。基于上述特征重要度的定义, 正向贪心搜索算法能被用来迭代地选择邻域决策系统中最重要特征子集。形式化地, 正向特征选择算法如下。

算法 2 基于邻域粗糙集的特征选择算法

输入: 邻域决策系统 $\langle U, C \cup D, V, f \rangle$, 邻域半径 δ_1

输出: 选择的特征子集 B

1. $B \leftarrow \emptyset$;
2. $\forall a_i \in C$, 计算 $Sig(a_i, B, D)$, 其中 $\gamma(\emptyset, D) = 0$;
3. 选择特征 a_j 按照:

$$Sig(a_j, B, D) = \max\{Sig(a_i, B, D) : \forall a_i \in C\};$$

4. 当 $Sig(a_j, B, D) > 0$ 时

4.1 $B \leftarrow B \cup a_j$;

4.2 $\forall a_i \in C - B$, 计算 $Sig(a_i, B, D)$;

4.3 选择特征 a_j 按照:

$$Sig(a_j, B, D) = \max\{Sig(a_i, B, D) : \forall a_i \in C - B\};$$

5. 输出特征子集 B 。

算法 2 把最快速增加邻域依赖度的特征依次加入到被选特征集合中, 直到邻域依赖度不再随特征的增加而增大为止。算法 2 的时间复杂度为 $O(|U|^2 |C|^2)$ 。

3.2 邻域协同分类器

为了缓解邻域分类器仅仅能够适用于较小信息粒下的分类的不足, 本节将采用协同表达的方式来替代多数投票原则, 旨在达到提升较大信息粒下分类器分类性能的目的, 提出的邻域协同分类器设计如下。

算法 3 邻域协同分类器(NCC)

输入: 训练集 X , 未知样本 y , 邻域半径 δ_2 , 正则化参数 λ , 单位矩阵 I

输出: 预测的 y 的类别

1. 对 X 的每一列采用 l_2 -范式进行归一化;
2. 对所有 $x_i \in X$, 计算距离 $\Delta(y, x_i)$;
3. 找出 y 的邻域 δ_2 内的样本集 $X_{\delta(y)}$;
4. 用 $X_{\delta(y)}$ 对 y 进行编码:

$$(\hat{\alpha}) = (X_{\delta(y)}^T X_{\delta(y)} + \lambda \cdot I)^{-1} X_{\delta(y)}^T y;$$

5. 对每类样本计算正则化残差:

$$r_i(y) = \frac{\|y - X_{\delta(y)i} \cdot \hat{\alpha}_i\|_2}{\|\hat{\alpha}_i\|_2};$$

6. 将最小 $r_i(y)$ 所在的类别 i 作为 y 的类别。

NCC 的基本流程是: 1) 找出未知样本的邻域; 2) 在邻域空间内计算未知样本的协同表达系数的解; 3) 对邻域空间内的每类训练样本分别计算重构误差; 4) 选择具有最小重构误差的类别作为未知样本的类别标记。

4 实验分析

4.1 数据集

为了验证邻域协同分类器的有效性,本文从UCI机器学习知识库选取了4组真实数据集进行实验分析,这4组数据集的一些基本性质如表1所列。

表1 实验数据集的基本性质

数据集	样本数	特征数	特征类型
Glass	214	9	连续型
Ionosphere	351	34	混合型
Parkinson	1208	26	连续型
Sonar	208	60	连续型

4.2 参数设置

本次实验将NCC与NC,CRC进行比较,实验选用配备2.7GHz处理器、8.00GB内存的PC机,编程环境为Matlab 2014b;在特征选择部分,为了获得较好的分类性能,根据专家经验把控制邻域信息粒大小的阈值 δ_1 分别设定为0.005,0.01,0.10和0.10;对于NCC和CRC,正则化参数 λ 被设定为0.01;此外,为了避免没有任何样本落入未知样本的邻域空间内的情形,NCC和NC中的邻域半径 δ_2 采用动态调节的方式,即有 $\delta_2 = \min(\Delta(y, x_i)) + \omega \cdot (\max(\Delta(y, x_i)) - \min(\Delta(y, x_i)))$,且 ω 的取值范围为0~1。

4.3 实验结果与讨论

采用分层随机抽样的方法对不同算法的分类性能进行评估。根据数据集中决策特征的个数将所有样本划分成一些独立的子样本集合,然后分别在每个子样本集合中随机选择60%的样本作为训练集,剩余40%的样本作为测试集。图2—图5为NCC,NC和CRC在4组数据集上分类性能的对比情况,其中带空心圆的线代表NC,带叉的线代表CRC,带实心三角形的线代表NCC。

从图2—图5可以发现,NC的分类准确率随着邻域信息粒尺寸的增大呈现明显的下降趋势。当邻域半径处于一个较小的区间时,邻域空间内的样本能够反映未知样本的局部信息,NC采用的多数投票原则对于判断未知样本的决策类别较为适用,为此,NC在较小的邻域半径下获得了良好的分类效果。然而,NCC的分类准确率大体上呈现出先保持增长、后保持稳定的趋势。当邻域信息粒尺寸较小时,邻域空间内的样本对于NCC而言不足以对未知样本进行重构,随着邻域半径的增大,越来越多的样本落入备选字典集合中,字典中样本数量的增加在一定程度上对NCC分类准确率的提升起到了促进作用。以Sonar数据集为例,当 $\omega < 0.4$ 时,NC的分类准确率优于NCC;而当 $\omega > 0.4$ 时,NCC的分类效果相比NC越来越有优势。此外,NCC通过局部训练样本对未知样本进行协同表达,大多数情况下能够在保持良好的分类精度的同时,很大程度地降低字典的规模。以Parkinson数据集为例, ω 从0.05开始增长,NCC的分类准确率开始优于CRC或与CRC相当,而此时字典样本仅占所有训练样本的极小一部分,由此可见,邻域协同表达是一种降低协同分类器计算复杂度的有效手段。

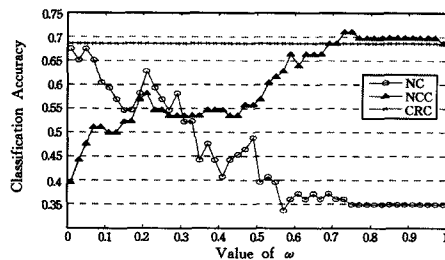


图2 Glass数据集上的分类准确率

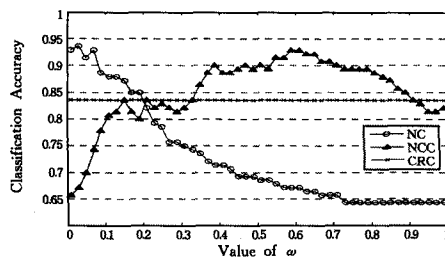


图3 Ionosphere数据集上的分类准确率

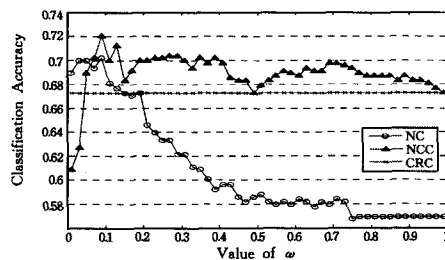


图4 Parkinson数据集上的分类准确率

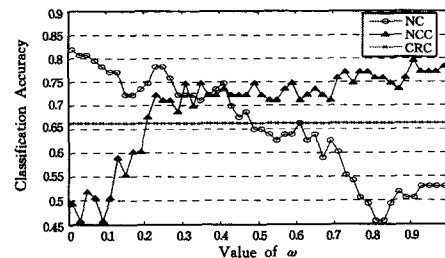


图5 Sonar数据集上的分类准确率

结束语 由于大尺寸信息粒下样本间的相似度过低,基于多数投票原则的邻域分类器并不适用于较大信息粒下的分类问题,为了缓解该问题,本文提出了一种基于邻域协同表达的分类方法,即邻域协同分类器(NCC)。NCC首先利用邻域粗糙集模型对分类任务进行特征选择,然后找出未知样本的邻域空间,最后在邻域空间内采用协同表达的方式对未知样本的类别进行判断。实验结果表明,与NC相比,NCC能够在较大尺寸信息粒下获得较为满意的分类效果;此外,NCC能够在保持良好分类精度的同时很大程度地缩减CRC中字典样本的规模,进而提升了分类的效率。

参考文献

- [1] PAWLAK Z. Rough sets; Theoretical aspects of reasoning about data [M]. Kluwer Academic Publishers, 1991.
- [2] PAWLAK Z, SKOWRON A. Rudiments of rough sets [J]. Information Sciences, 2007, 177(1): 3-27.
- [3] KANEIWA K, KUDO Y. A sequential pattern mining algorithm

- using rough set theory[J]. *International Journal of Approximate Reasoning*, 2011, 52(6): 881-893.
- [4] MARTINEZ I G, PEREZ R E B. Making decision in case-based systems using probabilities and rough sets[J]. *Knowledge-Based Systems*, 2003, 16(4): 205-213.
- [5] YANG X B, YAN X, XU S P, et al. New heuristic attribute reduction algorithm based on sample selection[J]. *Computer Science*, 2016, 43(1): 40-43. (in Chinese)
杨习贝, 颜旭, 徐苏平, 等. 基于样本选择的启发式属性约简方法研究[J]. *计算机科学*, 2016, 43(1): 40-43.
- [6] YANG X B, QI Y S, SONG X N, et al. Test cost sensitive multi-granulation rough set: Model and minimal cost selection [J]. *Information Sciences*, 2013, 250(11): 184-199.
- [7] HU Q H, PEDRYCZ W, YU D R, et al. Selecting discrete and continuous features based on neighborhood decision error minimization [J]. *IEEE Transactions on Systems Man & Cybernetics-Part B*, 2010, 40(1): 137-150.
- [8] HU Q H, YU D R, XIE Z X. Neighborhood classifiers [J]. *Expert Systems with Applications*, 2008, 34(2): 866-876.
- [9] HU Q H, YU D R, XIE Z X. Numerical attribute reduction based on neighborhood granulation and rough approximation [J]. *Journal of Software*, 2008, 19(3): 640-649. (in Chinese)
胡清华, 于达仁, 谢宗霞. 基于邻域粒化和粗糙逼近的数值属性约简[J]. *软件学报*, 2008, 19(3): 640-649.
- [10] WRIGHT J, YANGA Y, GANESH A, et al. Robust face recognition via adaptive sparse representation [J]. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2009, 31(2): 210-227.
- [11] ZHANG L, YANG M, FENG X C. Sparse representation or collaborative representation: Which helps face recognition? [C]// *Proceedings of the 2011 International Conference on Computer Vision*. IEEE Computer Society, 2011: 471-478.
- [12] WILSON D R, MARTINEZ T R. Improved heterogeneous distance functions [J]. *Journal of Artificial Intelligence Research*, 1997, 6(5): 1-34.
- [13] XU S P, YANG X B, SONG X N, et al. Prediction of protein structural classes by decreasing nearest neighbor error rate [C]// *Proceedings of the 2015 International Conference on Machine Learning and Cybernetics*. IEEE Computer Society, 2015: 7-13.
- [14] XU S P, YANG X B, YU H L, et al. Multi-label learning with label-specific feature reduction [J]. *Knowledge-Based Systems*, 2016, 104: 52-61.
-
- (上接第 233 页)
- [4] LIU L, XIONG X P. Least cache value replacement algorithm [J]. *Journal of Computer Applications*, 2013, 33(4): 1018-1022. (in Chinese)
刘磊, 熊小鹏. 最小驻留价值缓存替换算法[J]. *计算机应用*, 2013, 33(4): 1018-1022.
- [5] GAO M, WANG N H, LI D, et al. Data pre-fetching and caching algorithm based on templates [J]. *Application Research of Computers*, 2014, 31(11): 3240-3242, 3246. (in Chinese)
高萌, 王霓虹, 李丹, 等. 一种基于模板的数据预取和缓存算法[J]. *计算机应用研究*, 2014, 31(11): 3240-3242, 3246.
- [6] HEFEEDA M, NOORIZADEH B. On the benefits of cooperative proxy caching for peer-to-peer traffic [J]. *IEEE Transactions on Parallel and Distributed Systems*, 2010, 21(7): 998-1010.
- [7] LIANG W, BAYHAN S, KANGASHARJU J. Effects of cooperation policy and network topology on performance of in-network caching [J]. *IEEE Communications Letters*, 2014, 18(4): 680-683.
- [8] WU J L, YANG Q. A Web cache replacement algorithm based on collaborative filtering [J]. *Computer Engineering & Science*, 2015, 37(11): 2128-2133. (in Chinese)
吴俊龙, 杨清. 基于协同过滤的 Web 缓存替换算法研究[J]. *计算机工程与科学*, 2015, 37(11): 2128-2133.
- [9] HU Y Q, LI X N. A new proxy cache replacement mechanism of multimedia streams based on recommendation [J]. *Journal of Yanshan University*, 2015, 39(2): 139-144, 151. (in Chinese)
胡玉琦, 李晓娜. 一种新的基于推荐的流媒体代理缓存替换机制[J]. *燕山大学学报*, 2015, 39(2): 139-144, 151.
- [10] CAO M, LIU W Z. Research on cache replacement model based on multi-request mode under Hybrid architecture model [J]. *Computer Science*, 2015, 42(6): 175-180. (in Chinese)
曹旻, 刘文. 混合架构下多请求模式的缓存替换算法模型研究[J]. *计算机科学*, 2015, 42(6): 175-180.
- [11] LIN M W, YAO Z Q, XIONG J B. History-aware page replacement algorithm for NAND flash-based consumer electronics [J]. *IEEE Transactions on Consumer Electronics*, 2016, 62(1): 23-29.
- [12] WAN S G, HE X B, HUANG J Z, et al. An efficient penalty-aware cache to improve the performance of parity-based disk arrays under faulty conditions [J]. *IEEE Transactions on Parallel and Distributed Systems*, 2013, 24(8): 1500-1513.
- [13] HSIEH J W, KUAN Y H. DCCS: Double circular caching scheme for DRAM/PRAM Hybrid cache [J]. *IEEE Transactions on Computers*, 2015, 64(11): 3115-3127.
- [14] XIE R L. Research and implementation of the cache system In distributed search engine [D]. Xi'an: Northwest University, 2009. (in Chinese)
谢瑞莲. 分布式搜索引擎中缓存系统的研究与实现[D]. 西安: 西北大学, 2009.
- [15] WANG W J. Research on web cache and prefetching model based on access path mining [D]. Chengdu: Southwest Jiaotong University, 2011. (in Chinese)
王文建. 基于访问路径挖掘的 Web 缓存与预取模型研究[D]. 成都: 西南交通大学, 2011.
- [16] YANG L, DONG H Q, LIU G L. Current progress of caching techniques in storage [J]. *Microcomputer Applications*, 2015, 4(5): 1-9. (in Chinese)
杨琳, 董欢庆, 刘国良. 存储领域缓存技术的现状[J]. *网络新媒体技术*, 2015, 4(5): 1-9.
- [17] CHEN N J, LIN P. An user access feature driven semantic cache replacement policy in middleware [J]. *Journal of Guangxi University (Nat Sci Ed)*, 2010, 35(5): 787-792. (in Chinese)
陈宁江, 林盘. 用户访问特征驱动的中间件语义缓存替换策略[J]. *广西大学学报(自然科学版)*, 2010, 35(5): 787-792.
- [18] MA H Y, WANG B. Query results caching and prefetching in web search engines based on user characteristics [J]. *Journal of Chinese Information Processing*, 2012, 26(6): 19-26. (in Chinese)
马宏远, 王斌. 基于用户特性的搜索引擎查询结果缓存与预取[J]. *中文信息学报*, 2012, 26(6): 19-26.