

基于粗糙集理论的属性分类

邓九英^{1,2}, 毛宗源², 徐 宁¹

(1. 广东教育学院 计算机科学系, 广东 广州 510303; 2. 华南理工大学 自动化科学与工程学院, 广东 广州 510640)

摘要: 依据粗糙集理论, 分析决策表中条件属性的分类的变化, 使得决策子集和决策规则支持度发生变化的情况, 经过归纳推理得出选取最佳属性分类法的判定方法及算法. 最后, 通过仿真结果验证了算法2的有效.

关键词: 属性分类1; 支持子集2; 决策规则3; 支持度4

中图分类号: TP273 **文献标识码:** A

Attribute classification based on the rough sets theory

DENG Jiu-ying^{1,2}, MAO Zong-yuan², XU Ning¹

(1. Department of Computer Science, Guangdong Institute of Education, Guangzhou Guangdong 510640, China;

2. College of Automation Science and Technology, South China University of Technology, Guangzhou Guangdong 510640, China)

Abstract: According to Rough sets theory, it is analyzed that the variation of the condition attribute partition contributes to the support degree of decision subsets and decision rules in the decision system. The method of selecting the best attribute classification is induced. Criteria and algorithms for the best attribute classification are introduced. Finally, the validity of algorithm 2 is verified through a simulation result.

Key words: attribute classification 1; support subsets 2; decision rules 3; support degree 4

1 引言(Introduction)

粗糙集理论(rough sets theory, 简称RST)自20世纪80年代由Zdzislaw Pawlak提出后, 很快被人们接受并成功地应用于模糊数据的分析, RST是一种用于分析不精确、模糊、不确定数据的新型数学方法^[1]. RST是建立在分类机制的基础上, 将分类理解为在特定空间上的等价关系, 而等价关系构成了对该空间的划分^[2].

应用RST的分类数据能力, 把一个信息系统(U, A)转换成一个计算处理方便的数据关系(二维表), 对表中属性 $a \in A$ 的连续值做离散化处理, 人们对选取区间大小与分区边界的问题做了大量的研究工作^[3~6]. 但是, 在确定了离散化方法与分区边界处理原则的情况下, 使用同一种算法对表中的属性分类时, 对数据区间再划分(细分)会影响决策表的分析数据^[7,8]. 对于每一个属性的分类方法不同, 属性的取值及其等价类就不同, 由此直接导致对信息系统的属性约简结果会不同、导致知识获取后决策规则会不同. 为了选择一种合理的属性分类方法, 下面通过分析决策规则支持度的变化, 探求高效的属性分类方案, 并通过仿真试验对算法进行验证.

2 属性分类与支持子集(Attribute classification and support subset)

一个信息系统 S 是(U, A), 其中 $U = u_1, u_2, \dots, u_{|U|}$ 是有限非空集, 称为论域或对象空间, U 中的元素称为对象; $A = a_1, a_2, \dots, a_{|A|}$ 也是一个有限非空集, A 中的元素称为属性; 对于每个 $a \in A$, 有一个映射 $a : U \rightarrow a(U); a(U) = \{a(u) | u \in U\}$ 称为属性 a 的值域^[2,9].

一个信息系统可以用一个信息表来表示, 为了方便假设信息表中没有重复元组, 这时信息表是一个关系. 如果 $A = C \cup D, C \cap D = \emptyset$, 则称信息系统(U, A)为一个决策表, 其中 C 中的属性称为条件属性, D 中的属性称为决策属性.

2.1 属性的分类(Classification of attributes)

设(U, A)是一个信息系统, $U = u_1, u_2, \dots, u_{|U|}$, 对于每一个属性 $a \in A$, 引入一个 U 中的划分 U/a : 两个对象 $u, z \in U$ 在同一类中当且仅当 $a(u) = a(z)$. 可采用算法P对 U/a 进行分类.

算法 P 1) 初始化: $s = 1, i = 1, j = 1, a(u)$ (其中 $a \in A$), $V_s = u_i$;

2) $i = i + 1$, 循环比较 $a(V_j) = a(u_i)(j =$

$1, \dots, s)$;

3) 如果条件成立, $V_j = V_j \cup u_i$ (即令 u_i 在 V_j 中); 否则, $s+1 \rightarrow s, V_s = u_i$;

4) 重复步骤2) 3), 直至 $i = |U|$.

当算法结束时, 完成了对属性的划分 $U/a = V_1, V_2, \dots, V_s$.

2.2 支持子集(Support subset)

设 $W \in U$, 对于分类 U/a , 引入有关 W 的计算公式.

定义 1 W 的下近似为: $W^{(U/a)^-} = S_a(W) = \bigcup_{V \in U/a, V \subseteq W} V$; 子集 $S_a(W)$ 又称为关于属性 a 的支持子集(support subset).

定义 2 W 关于属性 a 的支持度(support degree)为: $spt_a(W) = |S_a(W)|/|U|$.

定义 3 W 的上近似为

$$W^{(U/a)^+} = \bigcup_{V \in U/a, V \cap W \neq \emptyset} V.$$

定义 4 W 关于属性 a 的近似精度为

$$acc_a(W) = |W^{(U/a)^-}|/|W^{(U/a)^+}|.$$

2.3 决策属性的支持度(Support degree of decision attribute)

设 $y \in D$ 是决策表 (U, A) 中的一个决策属性, $A = C \cup D, C \cap D = \emptyset$. 设 $U/y = W_1, W_2, \dots, W_t, 1 \leq t \leq |U|$. 考虑决策属性 $y \in D$ 关于 $a \in C$ 的支持子集.

推论 1 y 关于 $a \in C$ 的支持子集为: $S_a(y) = \bigcup_{W \in U/y} W^{(U/a)^-} = \bigcup_{W \in U/y} (\bigcup_{V \in U/a, V \subseteq W} V)$; 则 y 关于 a 的支持度为: $spt_a(y) = |\bigcup_{W \in U/y} W^{(U/a)^-}|/|U|$.

2.4 多个条件的支持度(Multi-condition supporting)

设 $W \in U$, 考虑子集 W 关于多个条件属性集的支持度.

定义 5 对于条件属性集 $X \subseteq C$, W 关于 X 的支持子集: $S_X(W) = W^{(U/X)^-} = \bigcup_{V \in U/X, V \subseteq W} V$; 而 W 关于 X 的支持度为 $spt_X(W) = |W^{(U/X)^-}|/|U|$.

令 $Y \subseteq D$ 是 $(U, C \cup D)$ 中的决策属性子集, 考虑 Y 关于 $X \subseteq C$ 的支持度.

推论 2 Y 关于 X 的支持子集为: $S_X(Y) = \bigcup_{W \in U/Y} W^{(U/X)^-} = \bigcup_{W \in U/Y} (\bigcup_{V \in U/X, V \subseteq W} V)$; 而 Y 关于 X 的支持度为: $spt_X(Y) = |\bigcup_{W \in U/Y} W^{(U/X)^-}|/|U|$.

3 几种属性分类情况的分析(Analysis of some attribute classification)

3.1 决策规则的产生(Generation of decision rules)

在决策表中, 最重要的是决策规则的产生.

设 $S = (U, A, V, f)$ 是一个决策表, $A = C \cup D, C \cap D = \emptyset$, 其中 C 为条件属性集, D 为决策属性集. 令 X_i 和 Y_j 分别代表 C 与 D 中的各个等价类, $\text{des}(X_i)$ 表示对等价类 X_i 的描述, 即等价类 X_i 对于各条件属性集的特定取值; $\text{des}(Y_j)$ 表示对等价类的描述, 即等价类 Y_j 对于各决策属性值的特定取值^[9]. 决策规则定义如下:

r_{ij} : $\text{des}(X_i) \rightarrow \text{des}(Y_j), Y_j \cap X_i \neq \emptyset$, 规则的确定性因子 $\mu(X_i, Y_j) = |Y_j \cap X_i|/|X_i|, 0 < \mu(X_i, Y_j) \leq 1$.

当 $\mu(X_i, Y_j) = 1$ 时, r_{ij} 是确定的; 当 $0 < \mu(X_i, Y_j) < 1$ 时, r_{ij} 是不确定的.

3.2 属性分类对决策规则支持度的影响(Variation of decision rule supporting infected by attribute classification)

在简化讨论问题时, 通常取信息系统中属性 $a \in A$ 分类数尽量少, 相应等价类 X_i 和 Y_j 的总数会减少, 对问题的求解会直观、更简单; 但生成的决策规则的准确性和可靠性会大受影响. 因此, 在数据预处理时, 应该取属性 $a \in A$ 的分类数尽量多, 相应等价类 X_i 和 Y_j 的总数会增多, 同时问题求解的算法会变得更复杂, 执行时间自然加长, 但可以换来更精确的决策规则, 应用价值提高.

属性 $a \in A$ 的值分类不是越多越好, 有时分类太多会使问题复杂化而又毫无意义; 这样, 如何选择属性分类的方法和属性值的划分, 直接关系到生成决策规则的可用价值, 下面通过具体例子的分析来寻求解决问题的途径.

表 1 决策表的两种属性分类

Table 1 Two attribute classifications of decision table

	分类1-分类2							
	u_1	u_2	u_3	u_4	u_5	u_6	u_7	u_8
U :								
C_1 :	1-1	1-2	1-2	0-0	0-0	0-1	0-1	0-0
C_2 :	0-0	1-1	2-2	0-0	1-1	2-2	1-1	2-2
D :	0-0	1-1	1-1	0-0	0-0	1-1	1-1	0-0

例 在表1的分类1中(“-”前的内容), 条件属性为 C_1 (头痛): 0/不痛, 1/痛; C_2 (体温): 0/正常, 1/高, 2/很高; 决策属性为 D (流感): 0/否, 1/是.

在表1的分类2中(“-”后的内容), 条件属性为 C_1 (头痛): 0/不痛, 1/痛, 2/很痛; C_2 (体温): 0/正常,

1/高, 2/很高; 决策属性为 D (流感): 0/否, 1/是。

把上面两种属性分类对应的信息系统(决策表)的计算数据列出, 如表2所示。

其中: $U/C_1 = V_{11}, V_{12}$ 为条件属性 C_1 的等价类; $U/C_2 = V_{21}, V_{22}, V_{23}$ 为条件属性 C_2 的等价类; U/C_1C_2 为条件属性 C_1 与 C_2 等价类的交集; $U/D =$

Y_1, Y_2 为决策属性 D 的等价类; $spt_{C_1}(D)$ 和 $spt_{C_2}(D)$ 分别为关于 C_1, C_2 对决策属性的支持度; $spt_{C_1C_2}(D)$ 为关于多条件属性的支持度。

根据表2中得出的支持子集和支持度, 计算决策表1两种分类的确定性规则与不确定性规则, 如表3所示。

表2 两种属性分类的支持子集和支持度

Table 2 Support subsets and support degree of two attribute classifications

	分类1的支持子集和支持度	分类2的支持子集和支持度
U/C_1 :	$u_1, u_2, u_3, u_4, u_5, u_6, u_7, u_8$	$u_2, u_3, u_1, u_6, u_7, u_4, u_5, u_8$
U/C_2 :	$u_1, u_4, u_2, u_5, u_7, u_3, u_6, u_8$	$u_1, u_4, u_2, u_5, u_7, u_3, u_6, u_8$
U/C_1C_2 :	$u_1, u_2, u_3, u_4, u_5, u_7, u_6, u_8$	$u_1, u_2, u_3, u_4, u_5, u_6, u_7, u_8$
U/D :	$u_1, u_4, u_5, u_8, u_2, u_3, u_6, u_7$	$u_1, u_4, u_5, u_8, u_2, u_3, u_6, u_7$
$spt_{C_1}(D)$:	$spt_{C_1}(Y_1) + spt_{C_1}(Y_2) = 0 + 0 = 0$	$spt_{C_1}(Y_1) + spt_{C_1}(Y_2) = (3 + 2)/8$
$spt_{C_2}(D)$:	$spt_{C_2}(Y_1) + spt_{C_2}(Y_2) = (2 + 0)/8$	$spt_{C_2}(Y_1) + spt_{C_2}(Y_2) = (2 + 0)/8$
$spt_{C_1C_2}(D)$:	$spt_{C_1C_2}(Y_1) + spt_{C_1C_2}(Y_2) = (2 + 2)/8$	$spt_{C_1C_2}(Y_1) + spt_{C_1C_2}(Y_2) = 8/8$

表3 两种属性分类的决策规则

Table 3 Decision rules of two attribute classifications

分类1的相应规则	分类2的相应规则
(头痛,1)且(体温,0)→(流感,0)*	(头痛,1)且(体温,0)→(流感,0)*
(头痛,1)且(体温,1)→(流感,1)*	(头痛,2)且(体温,1)→(流感,1)*
(头痛,1)且(体温,2)→(流感,1)*	(头痛,0)且(体温,0)→(流感,0)*
(头痛,0)且(体温,0)→(流感,0)*	(头痛,0)且(体温,1)→(流感,0)*
(头痛,0)且(体温,1)→(流感,1)	(头痛,2)且(体温,2)→(流感,1)*
(头痛,0)且(体温,1)→(流感,0)	(头痛,1)且(体温,2)→(流感,1)*
(头痛,0)且(体温,2)→(流感,1)	(头痛,1)且(体温,1)→(流感,1)*
(头痛,0)且(体温,2)→(流感,0)	(头痛,0)且(体温,2)→(流感,0)*

从表2和表3(用“*”标示的为确定性规则)可以看出, 决策属性的支持度高, 产生的决策规则的确定性就高。

4 属性分类法的判定(Judge according to attribute classification)

表1的分类2中的条件属性 C_1 或 C_2 , 还可以再划分为4个属性值: 0,1,2,3, 但相应的多条件支持度 $spt_{C_1C_2}(D)$ 不会变化(已经达到了最大值1); 生成决策规则的支持度不会变化(规则的支持度都已经为1)。相反, 条件属性值的再细分会使程序算法的复杂性增大、执行时间加长、效率降低; 会使决策规则的相似性增大、含义不明确。

根据上面的分析结果, 总结出满足支持度要求且具有最少属性值的属性分类法的判据和求解算法。

判据 1 已知决策表 $S = (U, A, V, f)$, $A = C \cup D, C \cap D = \emptyset, |U| = n$; 条件属性: $x_i \in C, i = 1, 2, \dots, |C|$; 决策属性: $Y \in D$ 。

对条件属性采用某种分类法得到的子类:
 $U/x_i = V_{i1}, V_{i2}, \dots, V_{im_i}, i = 1, 2, \dots, |C|, 1 \leq m_i \leq n$; 其中, m_i 的取值使决策规则的支持度最佳时, 应满足式子

$$m_i = \min |S_C(Y)| = |U|, Y \in D, i = 1, 2, \dots, |C|,$$

或者

$$m_i = \min |spt_C(Y)| = 1, Y \in D, i = 1, 2, \dots, |C|.$$

式中: $S_C(Y)$ (即 $S_{X_1, X_2, \dots, X_{|C|}}(Y), Y \in D$)为决策属性支持子集的交集; $spt_C(Y)$ (即 $spt_{X_1, X_2, \dots, X_{|C|}}(Y), Y \in D$)为决策属性的多条件支持度。

算法 1 采用逆向推理算法, 当 $m_i = n, i = 1, 2, \dots, |C|$ 时(一个属性值对应一个等价类), 对条件属性分类得到的子类都为: $U/x_i = u_1, u_2, \dots, u_n, i = 1, 2, \dots, |C|$; 则决策属性支持子集的交集也为: $S_C(Y) = u_1, u_2, \dots, u_n$, 显然, 决策属性的多条件支持度 $spt_C(Y) = 1$ 。

算法的主要步骤为

- 1) 设定初始值 $m_i = n, i = 1, 2, \dots, |C|$;
- 2) 按顺序取指针 $i = 1$, 循环做 $m_i - 1 \rightarrow m_i$, 其他 $m_j, j \neq i$ 值不变, 计算 $spt_C(Y)$ 的值, 直到 $m_i = h_i$ 时, $spt_C(Y) = 1$; 而 $m_i = h_i - 1$ 时, $spt_C(Y) < 1$;
- 3) 改变指针 $i + 1 \rightarrow i$, 重复步骤2), 直到 $i = |C|$, 停止.

当算法结束时, $H = h_1, h_2, \dots, h_{|C|}$ 为所求结果.

判据1适合较理想的情况, 算法1也是数学上的推导; 实际上, 在数据挖掘或决策控制领域涉及到的信息系统数据量很大, 很难达到判据1给出的标准, 这时可依据近似精度的判别方法.

判据 2 已知决策表 $S = (U, A, V, f)$, $A = C \cup D, C \cap D = \emptyset, |U| = n$; 条件属性: $x_i \in C, i = 1, 2, \dots, |C|$; 决策属性: $Y \in D$.

根据要求选取一组精度值: $K = k_i, 0 < k_i \leq 1, i = 1, 2, \dots, |C|$;

对条件属性采用某种分类法得到的子类: $U/x_i = V_{i1}, V_{i2}, \dots, V_{im_i}, i = 1, 2, \dots, |C|, 1 \leq m_i \leq n$;

其中, m_i 的取值使关于决策属性的条件属性分类满足近似精度的期望值: $acc_{x_i}(Y) \geq k_i, i = 1, 2, \dots, |C|$.

算法 2 采用正向推理算法, 分别取初始值 $m_i = 2, i = 1, 2, \dots, |C|$ (这时属性分类近似精度最低); 并逐步增加其值, 使属性分类的近似精度达到期望值.

算法的主要步骤为:

- 1) 给定条件属性分类的期望精度值: $K = k_i, 0 < k_i \leq 1, i = 1, 2, \dots, |C|$, 设定初始值 $m_i = 2, i = 1, 2, \dots, |C|$;
- 2) 按顺序取指针 $i = 1$, 循环计算 $acc_{x_i}(Y) = |Y^{(U/x_i)^-}|/|U^{(U/x_i)^+}|$ 的值, $m_i + 1 \rightarrow m_i$, 直到 $m_i = h_i$ 时, $acc_{x_i}(Y) \geq k_i$; 而 $m_i = h_i - 1$ 时, $acc_{x_i}(Y) < k_i$;
- 3) 改变指针 $i + 1 \rightarrow i$, 重复步骤2), 直到 $i = |C|$, 停止.

当算法结束时, $H = h_1, h_2, \dots, h_{|C|}$ 为所求结果.

5 仿真结果(Simulation results)

考虑到实用性, 选择判据2的算法进行仿真, 采用数据Iris(来源于加利福尼亚大学的机器学习数据中心)^[10], 如表4所示. 对条件属性进行离散化处理, 采用较简单的均分法离散取值为 $(0, 1, 2, 3, \dots)$. 其中: C_1 -sepal length; C_2 -sepal width; C_3 -petal length; C_4 -petal width; D -类别.

根据仿真软件得出近似精度 $acc_{C_i}(D)$ 与对应属性 C_i 的离散取值数 (N_i), 如表5所示. 对应判据2的分类结果, 较优的条件属性 X_i 离散取值个数 h_i 与近似精度 $acc_{C_i}(Y)$ 为 (其中 $i = 1, 2, 3, 4$): $H = 4, 5, 4, 5$; $acc(Y) = 1, 0.77, 1, 0.96$, 以及 (次优): $H = 3, 4, 3, 4$; $acc(Y) = 0.75, 0.59, 0.96, 0.7$.

以上的仿真结果依据的训练数据应与测试数据来自同源且相同规模 (采样个数), 但采样时间点可以不同, 此时仿真结果的应用才准确.

表4 Iris的部分示例数据
Table 4 Partial sample of iris' data

U	u_1	u_2	u_3	u_4	u_5	u_6	u_7	u_8	u_9	u_{10}	u_{11}	u_{12}	u_{13}	u_{14}	u_{15}	u_{16}	u_{17}	u_{18}	u_{19}	u_{20}	u_{21}	u_{22}	u_{23}	u_{24}	u_{25}	u_{26}	u_{27}
C_1	4.9	4.6	5.4	5	4.8	5.4	5	5.2	5.1	6.4	6.6	6.1	6.7	5.6	6.3	5.7	6	6.3	5.8	6.7	6.4	6	6.2	6.1	6.4	5.8	6.7
C_2	3	3.1	3.9	3.4	3	3.9	3	3.5	3.8	3.2	2.9	2.9	3.1	2.5	2.5	2.6	2.7	2.3	2.7	2.5	2.7	2.2	2.8	3	3.1	2.7	3.3
C_3	1.4	1.5	1.7	1.5	1.4	1.3	1.6	1.5	1.9	4.5	4.6	4.7	4.4	3.9	4.9	3.5	5.1	4.4	5.1	5.8	5.3	5	4.8	4.9	5.5	5.1	5.7
C_4	0.2	0.2	0.4	0.2	0.1	0.4	0.2	0.4	1.5	1.3	1.4	1.4	1.1	1.5	1	1.6	1.3	1.9	1.8	1.9	1.5	1.8	1.8	1.8	1.8	1.9	2.5
D	0	0	0	0	0	0	0	0	0	1	1	1	1	1	1	1	1	1	2	2	2	2	2	2	2	2	2

表5 仿真结果 N_i 和 $acc_{C_i}(D)$
Table 5 Simulation results, N_i and $acc_{C_i}(D)$

$C_1 : N, acc(D)$	$C_2 : N, acc(D)$	$C_3 : N, acc(D)$	$C_4 : N, acc(D)$
$4 \rightarrow 1 \sim 0.75$	$5 \rightarrow 0.77 \sim 0.6$	$4 \rightarrow 1 \sim 0.97$	$5 \rightarrow 0.96 \sim 0.71$
$3 \rightarrow 0.74 \sim 0.38$	$4 \rightarrow 0.59 \sim 0.34$	$3 \rightarrow 0.96 \sim 0.04$	$4 \rightarrow 0.7 \sim 0.54$
$2 \rightarrow 0.37$ 以下	$2 \rightarrow 0.33$ 以下	$2 \rightarrow 0.03$ 以下	$3 \rightarrow 0.55 \sim 0.41$
			$2 \rightarrow 0.4$ 以下

6 结束语(Conclusion)

从仿真结果可以看到, 本文提出的算法能直观地反映出属性分类(取值个数)的不同对决策表参数的影响, 并能依据判据快速地由精度值确定属性的分类。

本文讨论的关系表已经作了属性约简等数据的预处理, 决策表中的条件属性的重要性大于0才有意义, 仿真结果仅考虑了条件属性的取值个数(分区数)变化引起分类精度改变的情况。如果信息系统的属性还没有做属性约简, 属性分类的方法不同还会影响到属性的重要性, 近似精度对属性分类多少的敏感度又是属性约简的重要依据, 可能导致不同的属性约简结果, 这方面的研究可参考文献[11,12]。

参考文献(References):

- [1] ZDZISLAW P. Rough sets[C] // *Proceedings of the ACM 23rd Annual Conference on Computer Science of the 1995*. New York: ACM Press, 1995, 2: 262 – 263.
- [2] ZDZISLAW P. *Rough Sets: Theoretical Aspects of Reasoning about Data*[M]. London: Kluwer Academic Publishers, 1991, 1: 9 – 49.
- [3] DALE E N, JANUSZ A. High range resolution radar signal classification: a partitioned rough set approach[C] // *Proceedings of the 33rd Southeastern Symposium on System Theory*. Athens: IEEE Press, 2001, 3: 21 – 24.
- [4] JITENDER S D, CHOUBEY S K, VIJAY V. Raghavan and Hayri sever. feature selection and effective classifiers[J]. *Journal of the American Society for Information Science*, 1998, 4(5): 423 – 434.
- [5] ROMAN S. Rough sets with strict and weak indiscernibility relations[J]. *International Journal of Intelligent Systems*, 2002, 2(2): 153 – 171.
- [6] PAWAN L, BUTZ C. Interval set classifiers using support vector machines[C] // *Proceedings of IEEE Annual Meeting on Fuzzy Information, Processing NAFIPS'04*. Banff, Alberta, Canada: IEEE Press, 2004, 6: 707 – 710.
- [7] LI Ye, HU Z H, CAI Y Z, et al. Feature selection via modified RSBR for SVM classifiers[C] // *Proceedings of American Control Conference*. Portland, OR: IEEE Press, 2005, 6: 1455 – 1459.
- [8] WANG L W, ZHANG L, ZHANG M. A method of pattern classification based on RS and NCA[C] // *Proceedings of Machine Learning and Cybernetics*. Xi'an, China: IEEE Press, 2003, 11: 3090 – 3094.
- [9] 张文修, 吴伟志, 梁吉业. 粗糙集理论与方法[M]. 北京: 科学出版社, 2001, 7: 24 – 34.
(ZHANG Wenxiu, WU Weizhi, LIANG Jiyie. *Rough sets theory and methodology* Random Process[M]. Beijing: Science Press, 2001, 7: 24 – 34.)
- [10] 陈文伟, 黄金才, 赵新昱. 数据挖掘技术[M]. 北京: 北京工业大学出版社, 2002, 12: 42 – 62.
(CHEN Wenwei, HUANG Jingcai, ZHAO Xinyi. *Data Mining Technology—Random Process*[M]. Beijing: Beijing University of Technology Publisher, 2002, 12: 42 – 62.)
- [11] 邓九英, 毛宗源, 徐宁. 基于粗糙集变分区的属性约简[J]. 华南理工大学学报(自然科学版), 2006, 9(34): 50 – 54.
(DENG Jiuying, MAO Zongyuan, XU Ning. Attribute reduction using attributes partition variation of rough set[J]. *Journal of South China University of Technology (Natural Science Edition)*, 2006, 9(34): 50 – 54.)
- [12] 乔斌, 李玉榕, 将静坪. 粗糙集理论的分层递阶约简算法及其信息理论基础[J]. 控制理论与应用, 2004, 21(2): 195 – 199.
(QIAO Bin, LI Yuerong, JIANG Jingping. Hierarchical reduction approach of rough sets theory and its basis on the information theory[J]. *Control Theory & Applications*, 2004, 21(2): 195 – 199.)

作者简介:

邓九英 (1962—), 女, 副教授, 目前研究方向为智能控制、数据挖掘技术、计算机仿真技术, E-mail: djy1111@126.com;

毛宗源 (1936—), 男, 教授, 博士生导师, 目前研究方向为现代控制理论及应用、电力电子技术, E-mail: auzymao@scut.edu.cn;

徐宁 (1957—), 女, 副教授, 目前研究方向为模式识别与人工智能、仿真技术。