

基于稀疏学习的自适应近邻分类算法

程德波¹, 苏毅娟²⁺, 宗 鸣¹, 朱永华³

(1. 广西师范大学 计算机科学与信息工程学院, 广西 桂林 541004; 2. 广西师范学院 计算机与信息工程学院, 广西 南宁 530023; 3. 广西大学 计算机与电子信息学院, 广西 南宁 530004)

摘 要: 为解决 k-NN 算法中固定 k 的选定问题, 引入稀疏学习和重构技术用于最近邻分类, 通过数据驱动 (data-driven) 获得 k 值, 不需人为设定。由于样本之间存在相关性, 用训练样本重构所有测试样本, 生成重构系数矩阵, 用 l_1 -范数稀疏重构系数矩阵, 使每个测试样本用它邻域内最近的 k (不定值) 个训练本来重构, 解决 k-NN 算法对每个待分类样本都用同一个 k 值进行分类造成的分类不准确问题。UCI 数据集上的实验结果表明, 在分类时, 改良 k-NN 算法比经典 k-NN 算法效果要好。

关键词: 稀疏学习; 重构技术; 数据驱动; l_1 -范数; 邻域

中图法分类号: TP181 文献标识号: A 文章编号: 1000-7024 (2015) 07-1912-05

doi: 10.16208/j.issn1000-7024.2015.07.045

Self-adaptive neighbor classification based on sparse learning

CHENG De-bo¹, SU Yi-juan²⁺, ZONG Ming¹, ZHU Yong-hua³

(1. College of Computer Science and Information Technology, Guangxi Normal University, Guilin 541004, China;
2. College of Computer and Information Technology, Guangxi Teachers Education University, Nanning 530023, China;
3. College of Computer, Electronics and Information, Guangxi University, Nanning 530004, China)

Abstract: To deal with the problem that k-NN algorithm selects the fixed k, the sparse learning and reconstruction techniques for classification were used, so that k value was obtained through data-driven without artificial set. Due to the existence correlation between the samples, every test sample was used to reconstruct all the training samples, reconstruction coefficient matrix was generated. The l_1 -norm was used to penalize the objective function, so that each test sample used its neighborhood nearest k (a variable value) training samples to reconstruct, which solved the problem of inaccurate classification caused by k-NN algorithm using the fixed k value. Results of experiments on UCI datasets show that the improved k-NN algorithm is better than the classical k-NN algorithm in terms of classification effect.

Key words: sparse learning; reconstruction techniques; data-driven; l_1 -norm; neighborhood

0 引 言

通过研究分析^[1,2]发现 k-NN (k-nearest neighbor) 算法存在邻域大小 k 难以取定的缺陷, 文献 [3] 提出 $k = \sqrt{n}$ (数据集样本数 n 超过 100 时效果较好), 这种取法通常要经过大量的实验和分析才能确定合适的 k 值, 通常得到的

结果并不理想, 而且每一个待分类样本都用固定邻域大小 k 最近的已知样本构造的分类器来进行分类, 但在实际应用中每个样本的邻域是不同的, 因此每个测试样本应选取不同大小的邻域来对其进行分类。结合下面的例子来讨论 k-NN 算法存在的缺陷。

例如, 如图 1 为一个二类数据集。当 k=1 时, 待分类样

收稿日期: 2014-07-12; 修订日期: 2014-09-15

基金项目: 国家自然科学基金项目 (61170131、61263035、61363009); 国家 863 高技术研究发展计划基金项目 (2012AA011005); 国家 973 重点基础研究发展计划基金项目 (2013CB329404); 广西自然科学基金项目 (2012GXNSFGA060004); 广西高校科学技术研究重点基金项目 (2013ZD041); 广西研究生教育创新计划基金项目 (YCSZ2015095、YCZ2015096)

作者简介: 程德波 (1990-), 男, 江西丰城人, 硕士, 研究方向为数据挖掘、机器学习; +通讯作者: 苏毅娟 (1972-), 女, 广西桂林人, 副教授, 研究方向为机器学习和数据挖掘; 宗鸣 (1990-), 男, 江苏泰州人, 硕士, 研究方向为数据挖掘、机器学习; 朱永华 (1994-), 男, 广西桂林人, 本科, 研究方向为数据挖掘。E-mail: 7294835098@qq.com

本 1 将会被误判为第二类, 待分类样本 2 分类正确; 当 $k=3$ 时, 待分类样本 2 将会误判为第一类; 当 $k=7$ 时, 同样待分类样本 2 将会误判为第一类。根据分析可以发现待分类样本 1 属于第一类, 适合它的 k 值应该可以取 3、5、7 等; 待分类样本 2 属于第二类, 适合它的 k 值应该可以取 1、9、11 等。

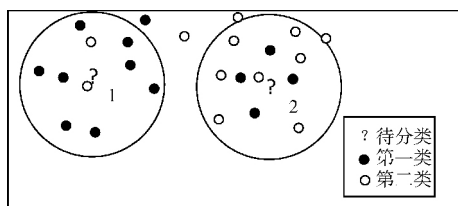


图 1 分类样本数据集

上述存在的问题显然是由用户定义 k 值产生, 即使是通过实验确定 k 值, 还是没有考虑数据样本之间的相关性, 不是数据驱动的。本文方法是数据驱动的, 由数据分布决定 k -NN 中 k (k 是学习而来)。具体的说, 本文借助稀疏学习^[4,5]和重构^[6]的相关知识用于 k -NN 分类来改良算法, 提出了一种 k -NN 分类算法——基于稀疏学习的自适应近邻分类算法 SL-SAN (self-adaptive neighbor classification based on sparse learning)。

1 SL-SAN 算法

1.1 重构技术与稀疏编码

重构技术: 重构是指用一组线性无关向量的线性组合来近似地表达一个给定的向量。待分类样本通过重构技术可以在一定半径的空间内找到与它近邻的已知样本。因此, 对于已知样本 $X = \{x_i\}_{i=1}^n \in R^{n \times d}$, n 为样本数, d 为维数; 待分类样本 $Y = \{y_i\}_{i=1}^m \in R^{m \times d}$, m 为样本数, d 为维数, 用已知样本 x_i 重构每一待分类样本 y_i 获取重构系数矩阵 $W \in R^{n \times m}$ 的目标函数如下

$$\arg \min_W \|w_i^T x_i - y_i\| \quad (1)$$

式中: w_i —— 重构系数矩阵 W 的列向量子数, 在文中解释为 x_i 与 y_i 的相似性度量。 w_i 值越大, 说明已知样本 y_i 与待分类样本 x_i 越相似, 但通过得到 w_i 还不能很容易的找到待分类样本大小为 k 个已知样本的邻域。以下介绍本文算法目标函数的导出。

最小二乘方模型可以用来处理回归问题也可以用来处理分类问题^[7], 本文采用最小二乘方来处理分类问题, 其模型如下

$$\|W^T X - Y\|_F^2 = \|\hat{Y} - Y\|_F^2 = \sum_{i=1}^n \sum_{j=1}^m (\hat{y}_{ij} - y_{ij})^2 \quad (2)$$

这里 $\|M\|_F^2 = \sum_{i=1}^n \sum_{j=1}^m m_{ij}^2$ (矩阵 M 全部元素的平方和) 为 Frobenius 范数, $W \in R^{n \times m}$ 是重构系数矩阵, $\hat{Y} = W^T X$ 。其中 $\|W^T X - Y\|_F^2$ 为重构误差, 由式 (2) 是凸的, 易知其

解为 $W^* = (XX^T)^{-1}XY$ 。在实际应用中 XX^T 不一定可逆, 然而正则化方法主要是为了改善反问题的不适应性, 通过构造与原问题相邻近的适应性问题来替代原不适应问题, 故本文在式 (2) 后加上 l_2 -范数正则化因子 $\|W\|_2^2$ 使其代替原小二乘方函数成为适应的可逆函数, 故目标函数转化为

$$\arg \min_W \|W^T X - Y\|_F^2 + \delta \|W\|_2^2 \quad (3)$$

其中, $\|W\|_2^2 = \sum_{i=1}^n \sum_{j=1}^m |w_{ij}|^2$, δ 为 l_2 -范式的正则化因子参数, 易知其解为: $W^* = (XX^T + \delta I)^{-1}XY$ 。由于稀疏学习能够很好捕捉到样本中的系数权重作为一种自然鉴别信息引入模型, 以此考虑样本数据分布结构信息和模型的鲁棒性, 且正则化因子 l_1 -范数已经被证明^[8,9]能使重构系数矩阵生成稀疏的系数矩阵, 比 l_2 -范数更容易导致稀疏性, 能保证本文 W 得到全局唯一最优解 (优化算法在下文算法 2 给出)。故本文将正则化因子 l_2 -范数替换成稀疏学习的正则化因子 l_1 -范数, 可得正则化目标函数的形式为

$$\arg \min_W \|W^T X - Y\|_F^2 + \lambda \|W\|_1 \quad (4)$$

其中, $\|W\|_1 = \sum_{i=1}^n \sum_{j=1}^m |w_{ij}|$, λ 为 l_1 -范数正则化因子参数, 通过增加 λ 的值可以使得 W 中的元素变成 0, 实现了特征子集的选择。

1.2 SL-SAN 分类算法

SL-SAN 算法将所有已知样本数据作为训练样本, 所有待分类样本作为测试样本; 用训练样本对测试样本进行重构得到重构系数矩阵; 借助稀疏学习, 采用 l_1 -范数正则化因子压缩重构系数矩阵, 使得重构系数矩阵成为稀疏矩阵。重构系数子数的稀疏位置不同确保了每个测试样本用不同的训练样本进行重构, 且稀疏约束下得到重构系数子数为 0 时, 表示该训练样本不参与该测试样本的重构, 因此不仅保留了样本之间的相关性, 而且获得了每个测试样本对应邻域不同个相似训练样本。算法主要思想描述如下:

假设测试样本 $X \in R^{n \times d}$, n 为样本数, d 为维数; 训练样本 $Y \in R^{m \times d}$, m 为样本数, d 为维数。根据重构和稀疏学习的知识采用式 (4) 函数: 用所有的训练样本 X 重构每一测试样本 Y , 得到重构系数矩阵并将其稀疏成为稀疏系数矩阵, 以此求解全局最优解 W 。式 (4) 即为 Lasso (the least absolute shrinkage and select operator)^[11], Lasso 是式 (2) 加 l_1 -范数而来, 并用 l_1 -范数惩罚目标函数使得绝对值较小的系数自动压缩为 0, 从而产生稀疏模型。并且正则化因子参数 λ 越大得到的稀疏密度越大, 即 W 为 0 的元素越多。

经式 (4) 最优化计算得到的 $W \in R^{n \times m}$, 在本文中 m 代表 m 个测试样本, n 代表 n 个训练样本跟每个测试样本之间的相似关系。且 w_{ij} 的值表示第 i 个测试样本与第 j 个训练样本的相似性。若 $w_{ij} > 0$, 说明第 i 个测试样本与第 j 个训练样本正相关, 且数值越大, 表明越相关; 若 $w_{ij} = 0$, 表示它们不相关; 若 $w_{ij} < 0$, 说明它们负相关。例如

$$W = \begin{bmatrix} 0.2 & 0 & 0.3 \\ 0 & 0.1 & 0 \\ 0 & 0 & 0.5 \\ 0.5 & 0 & 0.7 \end{bmatrix}$$

得到的最优化 W 是一个重构稀疏系数矩阵。 W 第 1 列不为 0 的数表明第 1 个测试样本跟第 1 个和第 4 个训练样本正相关, 因此, 第 1 个测试样本进行分类时设置 $k=2$; 同理, 第 2 个和第 3 个测试样本的 k 值依次应为 1 和 3。

学习 W 的过程是数据驱动 (data-driven), 因此能确保得到关系是最优, 并且通过 Lasso 的 l_1 范数使得每个测试样本需要的 k 值不同, 然后只需根据 W 取出与该测试样本相关的训练样本构造分类器, 进行分类。得到全局最优化结果 W 后, 可以对 Y 测试样本 (y 单个测试样本) 的每个测试样本构造分类器, 通过各分类器得到每个测试样本的类标签。

SL-SAN 算法伪代码描述如下所述:

算法 1: SL-SAN 算法

输入: X 训练样本, 并作规范化处理,

Y 测试样本 (y 单个测试样本);

λ 正则化参数;

输出: predict_label 预测标签;

all y do

{ (1) 求解如下最优化问题:

$$\arg \min_W \|W^T X - Y\|_F^2 + \lambda \|W\|_1$$

(2) 根据 (1) 优化得到的 W 获得 k

(3) 根据 (2) 得到每个测试样本对应的域大小 k

(4) 通过 k 值对每个 y 分类

(5) 求得 Y 的 predict_label

} end

SL-SAN 算法创新的使用 Lasso 来研究 k -NN 的 k 值固定问题, SL-SAN 算法虽然使用跟 Lasso 一样的目标函数, 但是与 Lasso 有两个区别: ① Lasso 一般用于分类或者回归等应用, 本文创新的利用 Lasso 介绍 k -NN 算法的 k 值固定值问题; ② 不同于常见使用 LARS (least angle regression) 方法求解 Lasso 的目标函数, 本文采用算法 2 来优化目标函数。SL-SAN 跟常见的 k -NN 算法比较: 首先, k -NN 是懒惰学习方法, 对每一个测试样本单独的求 k -NN 中的 k 值, SL-SAN 一次重构所有测试样本, 即考虑了测试样本的相关性也考虑了训练样本的相关性。其次, k -NN 算法采用用户自定义或者十折交叉验证方法获取 k 值, 且经常 k 值对所有测试例子一样, 而本文的 SL-SAN 通过 Lasso 重构方法能得到测试样本和训练样本之间的相关性矩阵 W 。经实验验证, SL-SAN 算法是优于经典 k -NN 方法。

1.3 算法优化

等式 (4) 是凸的, 但后面正则化项是非光滑的。故本

文提出一种有效的算法来光滑目标函数并求得最优解 W 。首先对 $w_i (1 \leq i \leq m)$ 求导并命其为 0, 可得

$$XX^T w_i - X y^{(i)} + \lambda D_i w_i = 0 \quad (5)$$

这里 $D_i (1 \leq i \leq m)$ 是对角矩阵, 第 k 个对角元素为

$$\frac{1}{2 \|w_{ki}\|}。所以$$

$$w_i = (XX^T + \lambda D_i)^{-1} X y^{(i)} \quad (6)$$

注意到 D_i 依赖于 W , 因此它们也是未知的。根据文献 [12] 本文提出采用一种迭代算法来求解最优解 W , 即算法 2 如下:

算法 2: 目标函数优化算法

输入: X, Y

输出: $W^{(t)} \in R^{n \times m}$

初始化 $W^1 \in R^{n \times m}$, $t=1$;

do {

(1) 计算对角矩阵 $D_i^{(t)} (1 \leq i \leq m)$, 这里 $D_i^{(t)}$ 为第 k

个对角元素为 $\frac{1}{2 \|w_{ki}^{(t)}\|}$

(2) For 每个 $i (1 \leq i \leq m)$

$$w_i^{(t+1)} = (XX^T + \lambda D_i^{(t)})^{-1} X y^{(i)}$$

(3) $t=t+1$;

} until 收敛

证明: 根据算法的第 2 步可得到

$$W^{(t+1)} = \min_W \text{Tr}(X^T W - Y)^T (X^T W - Y) + \lambda \sum_{i=1}^m w_i^T D_i^{(t)} w_i \quad (7)$$

因此有

$$\begin{aligned} & \text{Tr}(X^T W^{(t+1)} - Y)^T (X^T W^{(t+1)} - Y) + \lambda \sum_{i=1}^m (w_i^{(t+1)})^T D_i^{(t)} w_i^{(t+1)} \\ & \leq \text{Tr}(X^T W^{(t)} - Y)^T (X^T W^{(t)} - Y) + \lambda \sum_{i=1}^m (w_i^{(t)})^T D_i^{(t)} w_i^{(t)} \\ & \Rightarrow \text{Tr}(X^T W^{(t+1)} - Y)^T (X^T W^{(t+1)} - Y) + \\ & \quad \lambda \sum_{i=1}^d \sum_{j=1}^m \left(\frac{(w_{ij}^{(t+1)})^2}{2 \|w_{ij}^{(t)}\|} - \|w_{ij}^{(t+1)}\| + \|w_{ij}^{(t)}\| \right) \leq \\ & \quad \text{Tr}(X^T W^{(t)} - Y)^T (X^T W^{(t)} - Y) + \\ & \quad \lambda \sum_{i=1}^d \sum_{j=1}^m \left(\|w_{ij}^{(t)}\| + \frac{(w_{ij}^{(t)})^2}{2 \|w_{ij}^{(t+1)}\|} - \|w_{ij}^{(t)}\| \right) \\ & \Rightarrow \text{Tr}(X^T W^{(t+1)} - Y)^T (X^T W^{(t+1)} - Y) + \lambda \sum_{i=1}^d \sum_{j=1}^m \|w_{ij}^{(t+1)}\| \\ & \leq \text{Tr}(X^T W^{(t)} - Y)^T (X^T W^{(t)} - Y) + \lambda \sum_{i=1}^d \sum_{j=1}^m \|w_{ij}^{(t)}\| \end{aligned}$$

根据文献 [13] 对于任意向量 w 和 w_0 , 我们有 $\|w\| - \frac{\|w\|_2^2}{2 \|w_0\|_2} \leq \|w_0\|_2 - \frac{\|w_0\|_2^2}{2 \|w_0\|_2}$ 。因此最后一步成立, 即算法在每次迭代中减小目标值。

$W^{(t)}$, $D_i^{(t)} (1 \leq i \leq m)$ 在收敛处满足等式 (5)。由于问

题 (5) 是一个凸问题, 满足等式 (7) 意味着 W 对于问题 (5) 来说是一个全局最优解。因为在每一次迭代时都有封闭形式解, 因此, 算法 2 可以将问题 (5) 收敛到它的全局最优解, 故我们的算法收敛非常快。

2 实验结果与分析

实验的评价指标采用分类准确率, 分类准确率越大表明效果分类效果越好。计算公式如下

$$Accuracy = \frac{n_{correct}}{n} \tag{8}$$

式中: n ——元组数, $n_{correct}$ ——正确分类的元组数。

实验部分使用的数据集来源于 UCI^[13,14], 实验数据集介绍见表 1。

表 1 用于实验的数据集

Dataset	# of instances	# of features	# of classes
adult	1605	113	2
cleveland	214	13	2
sonar	208	60	2
seeds	210	7	3

本文算法用 MATLAB 语言编程实现, 并在 win7 系统下 MATLAB 7.1 软件上进行实验。实验使用 10 折交叉验证法, 从所有样本中取出一个作为测试样本 (作为待分类样本), 余下的作为训练样本。通过用表 1 数据集进行实验, 每个数据集重复实验 10 次, 并取 10 次实验准确率的均值加上方差可以得到表 2 实验结果。

表 2 SC-SAN、k-NN 准确率加上方差

Dataset	SC-SAN	k-NN
adult	0.8025±0.0004	0.7794±0.0007
cleveland	0.8381±0.0024	0.7619±0.0025
sonar	0.8500±0.0039	0.7650±0.0106
seeds	0.9381±0.0010	0.8762±0.0021

通过表 2 可以发现: 从准确率来看, 4 个数据集在使用算法 SL-SAN 分类时所得准确率比经典 k-NN 算法分类时要高, 且分类准确率提高了 3%~8%; 从方差的角度来看, SL-SAN 算法比 k-NN 算法具有很好的稳定性。通过图 2 可以更直观的看出 SL-SAN 算法的优越性。经实验证实 SL-SAN 算法比经典 k-NN 算法分类效果要好。

3 结束语

SL-SAN 算法创新使用稀疏学习和重构技术来介绍 k-NN 算法的 k 值固定值问题, 并解决了 k 值难以取定的难题, 完善了 k-NN 算法没有考虑待分类样本与所有已知样本之间相关性的缺陷。具体表现为, SL-SAN 算法学习得到 W 的过程是数据驱动过程, 通过 W 获得测试样本的 k 邻域自然也是自适应的, 无需人为设定, 并且重构后可以得知待分类样本分布情况, 得到的 k (不固定) 个已知样本是充分考虑了待分类样本与已知样本之间的相关性。本文算法适用于 k-NN 分类算法的适用范围。

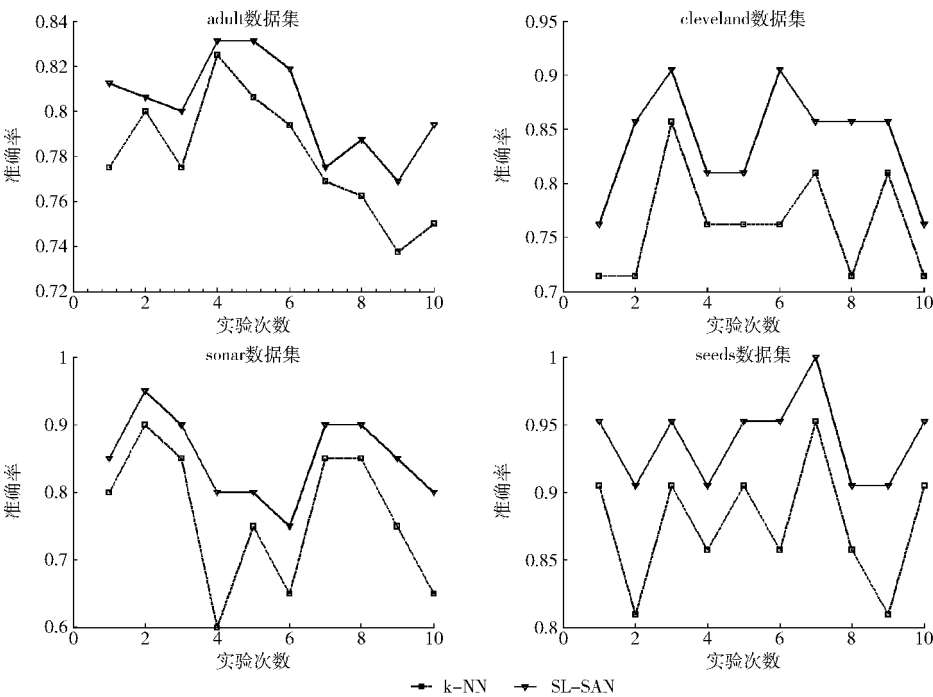


图 2 4 个数据集用两种算法分类 10 次实验准确率的对比

参考文献:

- [1] Zhang Shichao. Cost-sensitive classification with respect to waiting cost [J]. Knowledge-Based System, 2010, 23 (5): 369-378.
- [2] Zhu Xiaofeng, Zhang Shichao, Jin Zhi, et al. Missing value estimation for mixed-attribute data sets [J]. IEEE Transactions Knowledge Data Engineering, 2011, 23 (1): 110-121.
- [3] Lall U, Sharma A. A nearest-neighbor bootstrap for resampling hydrologic time series [J]. Water Resour Res, 1996, 32 (3): 679-693.
- [4] Jenatton R, Gramfort A, Michel V, et al. Mutual scale mining of fMRI data with hierarchical structured sparsity [J]. SIAM Journal on Imaging Sciences, 2012, 5 (3): 835-856.
- [5] Zhu Xiaofeng, Huang Zi, Shen Hengtao, et al. Dimensionality reduction by mixed-kernel canonical correlation analysis [J]. Pattern Recognition, 2012, 45 (8): 3003-3016.
- [6] KANG P, CHO S. Locally linear reconstruction for instance-based learning [J]. Pattern Recognition, 2008, 41 (11): 3507-3518.
- [7] LIANG Jinjin, WU De. Sparse least square support vector machine with L1 norm [J]. Compute Engineering and Design, 2014, 35 (1): 293-296 (in Chinese). [梁锦锦, 吴德. 稀疏 L1 范数最小二乘支持向量机 [J]. 计算机工程与设计, 2014, 35 (1): 293-296.]
- [8] Zhu Xiaofeng, Huang Zi. Sparse hashing for fast multimedia search [J]. ACM Transactions on Information System, 2013, 31 (2): 1-24.
- [9] Zhu Xiaofeng, Zhang Lei, Huang Zi. A sparse embedding and least variance encoding approach to hashing [J]. IEEE Transactions on Image Processing, 2014, 23 (9): 3737-3750.
- [10] Zhu Xiaofeng, Huang Zi, Cui Jiangtao, Hengtao Shen: Video-to-shot tag propagation by graphsparse group lasso [J]. IEEE Transactions on Multimedia, 2013, 15 (3): 633-646.
- [11] Zhu Xiaofeng, Huang Zi. Self-taught dimensionality reduction on the high-dimensional small-sized data [J]. Pattern Recognition, 2013, 46 (1): 215-229.
- [12] Zhu Xiaofeng, Huang Zi, Shen Hengtao, et al. Linear cross-modal hashing for efficient multimedia search [C] //ACM Multimedia, 2013: 143-152.
- [13] UCI repository of machine learning datasets [DB/OL]. [2014-05-06]. <http://archive.ics.uci.edu/ml/>, 2014.
- [14] LE SONG. Codes and data[EB/OL]. [2014-05-06]. <http://www.cc.gatech.edu/~lsong/code.htm>, 2014.

(上接第 1911 页)

参考文献:

- [1] LIU Qingjie, WANG Xiaoying, WANG Maofa. research on reasoning for fault diagnosis based on BP and ACO algorithms [J]. Computer Measurement&Control, 2012, 20 (6): 1460-1462 (in Chinese). [刘庆杰, 王小英, 王茂发. 基于 BP 算法和 ACO 算法的故障诊断推理研究 [J]. 计算机测量与控制 [J]. 2012, 20 (6): 1460-1462.]
- [2] DENG Ming, JIN Yezhuang. Aircraft engine fault diagnosis [M]. Beijing: Beihang University Press, 2011 (in Chinese). [邓明, 金业壮. 航空发动机故障诊断 [M]. 北京: 北京航空航天大学出版社, 2011.]
- [3] HU Chunling. Bayesian network structure learning and its application [D]. Anhui: Hefei University of Technology, 2011 (in Chinese). [胡春玲. 贝叶斯网络结构学习及其应用研究 [D]. 安徽: 合肥工业大学, 2011.]
- [4] ZHAO Jinbao, DENG Wei, WANG Jian. Bayesian network-based urban road traffic accidents analysis [J]. Journal of Southeast University (Natural Science Edition), 2011, 41 (6): 1301-1302 (in Chinese). [赵金宝, 邓卫, 王建. 基于贝叶斯网络的都市道路交通事故分析 [J]. 东南大学学报 (自然科学版), 2011, 41 (6): 1301-1302.]
- [5] Peng Y, Zhang S, Pan R. Bayesian network reasoning with uncertain evidences [J]. International Journal of Uncertainty Fuzziness and Knowledge-based Systems, 2010, 18 (5): 539-564.
- [6] Xu E, Fan L, Li S, et al. Research on preprocess approach for uncertain system based on rough set [C] //1st International Conference on Swarm Intelligence, China, 2010: 656-663.
- [7] Park H, Cho S. A modular design of Bayesian networks using expert knowledge: Context-aware home service robot [J]. Expert Systems with Applications, 2012, 39 (3): 2629-2642.
- [8] Reeve J, DeFreitas D, Sellares J. A Bayesian network for predicting diagnoses and risk of failure in kidney transplant biopsies suggests inaccuracy in the Banff system [J]. American Journal of Transplantation, 2011, 11 (2): 271-271.
- [9] Uhart M, Bourguignon Maire L, Ducher M. Bayesian networks as decision-making tools to help pharmacists evaluate and optimize hospital drug supply chain [J]. European Journal of Hospital Pharmacy-science and Practice, 2012, 19 (6): 519-524.
- [10] Glendining N, Pollino C. Development of Bayesian network decision support tools to support river rehabilitation works in the lower snowy river [J]. Human and Ecological Risk Assessment, 2012, 18 (1): 92-114.