

非线性相关的失效数据联合聚类分析与预测

卢 旭¹, 王慧强¹, 吕 晓², 冯光升¹, 林俊宇¹

(1. 哈尔滨工程大学 计算机科学与技术学院, 150001 哈尔滨, luxu_hrbeu@yahoo.cn;

2. 海军工程大学 电子工程学院, 430033 武汉)

摘 要: 为了实现对高性能计算机系统大规模失效数据的自动分析和预测, 引入了基于信息论的联合聚类思想提取非线性相关失效数据对象. 根据失效特征非线性相关性对失效数据进行归类, 并给出用于聚类的失效特征标签的定义, 以此为基础提出以互信息熵作为相似性度量的非线性相关失效数据联合聚类算法, 并从理论上论述了算法的收敛性和局部最优性. 实验结果显示联合聚类分析算法具有良好的运行性能, 成功聚类出非线性相关的失效数据对象, 验证了联合聚类分析方法用于失效预测的有效性.

关键词: 失效分析; 失效预测; 非线性相关性; 联合聚类

中图分类号: TP393.0

文献标志码: A

文章编号: 0367-6234(2011)03-0080-05

Nonlinearly correlated failure data Co-clustering analysis and prediction

LU Xu¹, WANG Hui-qiang¹, LÜ Xiao², FENG Guang-sheng¹, LIN Jun-yu¹

(1. College of Computer Science and Technology, Harbin Engineer University, 150001 Harbin, China, luxu_hrbeu@yahoo.cn;

2. College of Electronic Engineering, Naval Engineer University, 430033 Wuhan, China)

Abstract: To realize automatic analysis and prediction large-scale failure data for high-performance computer systems, an information-theoretic Co-clustering method for large scale nonlinearly correlated failure data was proposed. The failure data were sorted according to the nonlinear correlation of failure features and the failure feature signatures were defined. Then a nonlinearly correlated failure data Co-clustering algorithm was proposed and measured using mutual information entropy. Moreover, the convergence and local optimality of Co-clustering algorithm were proved theoretically. Experimental results on labeled failure data showed that the Co-clustering analysis algorithm outperformed other clustering analysis algorithms and the proposed Co-clustering analysis method had the features of rationality and effectiveness for discovering the nonlinearly correlated failure patterns and could help to predict underlying failures.

Key words: failure analysis; failure prediction; nonlinear correlation; Co-clustering

随着高性能计算机系统应用的日益普及和对高可用性的严格约束, 如何预测并避免系统可能的失效成为当前一个研究热点^[1-2]. 实现失效预测首先需要对系统所生成的大规模失效数据进行分析, 其研究关键就是将具有相似失效模式的失效样本聚集在一起, 为后续失效预测与快速恢复提供决策依据. 失效数据对象间的一个重要关系就是非线性相关性^[3], 可用以反映特征向量之

间的依赖性或者关联关系, 随着近年高性能计算机系统失效数据的对外开放^[4-5], 此类失效数据分析和预测问题的研究逐渐引起了广泛关注.

目前对系统失效数据的分析主要有: 1) 分析系统运行历史记录, 计算失效平均间隔时间; 2) 提取失效事件时序特征, 通过辨别系统失效模式预测失效. 这 2 类方法尽管在各自的系统上表现优异, 但没有有机融合失效特征进行综合评判, 或者对于冗余、重复数据需要人为进行筛选, 难以得到推广^[6-7]. 在分布式、高性能计算机系统中, 失效数据往往具有稀疏、高维等特征, 且有价值的非线性相关性主要存在于维度子空间内, 而此类

收稿日期: 2009-09-20.

基金项目: 国家高技术研究发展计划资助项目(2007AA01Z401);
国家自然科学基金资助项目(90718003, 60973027).

作者简介: 卢 旭(1983—), 男, 博士研究生;
王慧强(1960—), 男, 教授, 博士生导师.

单向聚类方法难以发现这些维度子空间内的局部相关性,增加了失效预测的不确定性.本文分析失效特征的非线性相关性,引入联合聚类思想(Co-clustering)^[8],以互信息熵损失差作为联合聚类度量标准,提出非线性相关失效数据联合聚类算法,实现非线性相关失效数据自动分析与预测.

1 失效预测基本模型

1.1 失效预测框架

尽管目前已经研究出多种容错机制,致力于避免故障发生后导致系统服务中断,但在实际应用中失效的发生总是不可避免,因此要保证系统的高可用性就必须为系统提供失效预测的能力.根据预测时间的长短,失效预测可分为长期预测和短期预测,其中:长期预测大多通过分析系统可靠性模型来获得失效事件的季节性分布;短期失效则通过失效事件的关联性来预见可能发生的失效,失效预测框架如图1所示.在 t 时刻预测在 $t + \Delta t$ 时间段内是否发生失效,其中: Δt 为预先时间; Δt_w 为 Δt 的下限表示的失效预恢复所需时间; Δt_p 为预测期,描述了整个预测时间窗口的长度.显然,预测时间窗口越长,则时间窗口内发生失效的可能性越大.

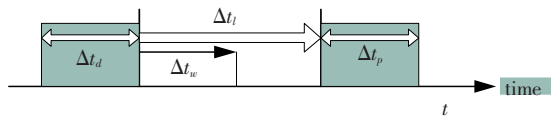


图1 失效预测基本框架

1.2 失效特征相关性分析

对大规模失效数据集的分析首先需要从运行系统中提取出能够反映与失效事件关联的系统状态特征.这些特征应当能够表示出系统在正常执行时与失效事件发生时的差异,同时还需要捕捉到不同失效事件中间的时空关联性.失效特征应该能够反映系统失效与正常运行下的本质区别,同时还需要能够反映失效事件的时空关联性.一般而言,特征空间可分为强相关特征、弱相关特征和无关特征^[9].设 F 是特征集合, f_i 是一个特征, $S_i = F - \{f_i\}$,特征相关性的形式化定义为:

定义1 特征 f_i 是强相关的当且仅当 $P(C|f_i, S_i) \neq P(C|S_i)$.

定义2 特征 f_i 是弱相关的当且仅当 $P(C|f_i, S_i) = P(C|S_i)$,且存在 $S'_i \subset S_i$,使得 $P(C|f_i, S'_i) \neq P(C|S'_i)$ 强相关特征是对类的分布构成影响的特征,弱相关特征则只在一定条件下影响不同类之间的分布,以此为基础可给出失效特征非线性相关性定义.

定义3 特征 f_i 是非线性相关的当且仅当特

征是强相关或弱相关,且 $\rho = 0$ 为

$$\rho = \frac{\sum_i (x_i - \bar{x}_i)(y_i - \bar{y}_i)}{\sqrt{\sum_i (x_i - \bar{x}_i)^2} \sqrt{\sum_i (y_i - \bar{y}_i)^2}} \quad (1)$$

式中 ρ 被称为线性相关系数且 $x, y \in F$.

通过比对系统各部分失效发生的关联性,给出用于失效预测特征提取的性能指标.这些指标可从系统事件日志中提取,例如处理器与存储空间利用率,通讯与输入输出操作的容量等.而固定时间窗口内的失效事件数、失效类型与平均失效间隔时间等则用来构建系统失效的统计模型.

定义4 失效特征标签由七元组给出定义: tuple(fID, time, fLoc, fType, usrUtil, pktCount, ioCount). 其中:fID, time, fLoc, fType 分别为失效事件编号、时间戳、失效定位和失效类型;usrUtil 为系统失效时资源利用率;pktCount, ioCount 分别为衡量特征提取阶段的数据包与通信请求数.

2 失效数据聚类方法

基于信息论的联合聚类方法从数据矩阵行维与列维2个方向上聚类,在数据分析、协同过滤等研究领域有着广泛应用^[10].假设 m 个待聚类失效样本,记作 $\{x_1, x_2, \dots, x_m\}$;若每个失效事件特征标签用 n 维特征向量 Y 描述,记作 $\{y_1, y_2, \dots, y_n\}$. $p(X, Y)$ 为随机变量 X 和 Y 的联合概率分布,由于 X 和 Y 都是离散型随机变量, $p(X, Y)$ 可以用一个 $m \times n$ 矩阵表示,矩阵中每个元素 $p(x, y)$ 表示失效事件 x 和失效特征标签 y 联合发生的概率.

定义5 联合聚类映射定义为

$$C_X: \{x_1, x_2, \dots, x_m\} \rightarrow \{\hat{x}_1, \hat{x}_2, \dots, \hat{x}_k\}, \quad (2)$$

$$C_Y: \{y_1, y_2, \dots, y_n\} \rightarrow \{\hat{y}_1, \hat{y}_2, \dots, \hat{y}_l\}.$$

式中:随机变量 $\hat{X}, \hat{Y}$ 分别为联合聚类后失效事件与特征标签集合; k, l 分别为联合聚类后失效事件簇和失效特征簇的数量.

联合聚类映射可简写为 $\hat{X} = C_X(X)$ 和 $\hat{Y} = C_Y(Y)$.文献[8]引入了互信息熵(mutual information)的概念对联合聚类问题进行量化分析,并给出了联合聚类映射函数的最优解形式.

定义6 最优联合聚类应满足聚类前后的互信息熵差最小化计算为

$$\operatorname{argmin} \{I(X; Y) - I(\hat{X}; \hat{Y})\}. \quad (3)$$

其中:互信息熵计算为

$$I(X; Y) = \sum_{\hat{x}} \sum_{\hat{y}} \sum_{x \in \hat{x}} \sum_{y \in \hat{y}} p(x, y) \log \frac{p(x, y)}{p(x)p(y)}. \quad (4)$$

对于联合聚类映射 $\hat{X} = C_X(X)$ 和 $\hat{Y} = C_Y(Y)$,互信

息熵损失由 Kullback-Leibler 距离计算,计算公式为

$$I(X;Y) - I(\hat{X};\hat{Y}) = D(p(X,Y) \| q(X,Y)). \quad (5)$$

式中: $D(\cdot \| \cdot)$ 为 KL 散度,或称为相关熵; $q(x,y) = p(\hat{x},\hat{y})p(x|\hat{x})p(y|\hat{y})$. 在求解联合聚类算法的映射函数时, KL 散度度量被进一步表述为 2 种对称形式,用于实现行聚类和列聚类,计算公式分别为

$$D(p(X,Y) \| q(X,Y)) = \sum_{\hat{x}} \sum_{x: C_X^{(t)}(x) = \hat{x}} p(x) \cdot D(p(Y|x) \| q(Y|\hat{x})), \quad (6)$$

$$D(p(X,Y) \| q(X,Y)) = \sum_{\hat{y}} \sum_{y: C_Y^{(t)}(y) = \hat{y}} p(y) \cdot D(p(X|y) \| q(X|\hat{y})). \quad (7)$$

在上述分析基础上,提出非线性相关的失效数据联合聚类分析算法. 算法分 3 个部分: 1) 根据已知的失效事件与特征标签联合分布进行初始化,并给出初始化映射函数; 2) 按照互信息熵差计算方法分别进行行聚类和列聚类,最终获得新的联合聚类子集; 3) 判断联合聚类后的数据矩阵是否满足终止条件,如果不满足则继续进行聚类迭代直至满足收敛条件为止. 算法详细步骤如图 2 所示.

算法 1 非线性相关失效数据联合聚类算法

输入: $p(X,Y)$ // 二维联合分布; ε // 任意小整数; k, l // 联合聚类簇与列簇数.

输出: C_X 和 C_Y // 联合聚类映射目标函数.

1) 初始化映射函数: set $t = 0$, $C_Y = C_Y^0$, $C_X = C_X^0$, 计算

$$q^0(\hat{X}, \hat{Y}), q^0(X|\hat{X}), q^0(Y|\hat{Y});$$

2) 行聚类更新, 根据失效特征标签计算新的类别归属为

$$C_X^{(t+1)}(x) = \underset{\hat{x}}{\operatorname{argmin}} D(p(Y|x) \| q^{(t)}(Y|\hat{x}));$$

3) 计算行聚类后二维变量联合分布为 $q^{t+1}(\hat{X}, \hat{Y}), q^{t+1}(X|\hat{X})$,

$$q^{t+1}(Y|\hat{Y}), q^{t+1}(X|\hat{y}), 1 \leq \hat{y} \leq l;$$

4) 列聚类更新, 根据失效特征标签计算新的类别归属为

$$C_Y^{(t+2)}(y) = \underset{\hat{y}}{\operatorname{argmin}} D(p(X|y) \| q^{(t+1)}(X|\hat{y}));$$

5) 计算列聚类后二维变量联合分布为 $q^{t+2}(\hat{X}, \hat{Y}), q^{t+2}(X|\hat{X})$,

$$q^{t+2}(Y|\hat{Y}), q^{t+2}(Y|\hat{x}), 1 \leq \hat{x} \leq k;$$

6) 判断联合聚类是否达到局部最优, 若达到则算法终止: if

$$D(p(X,Y) \| q^{(t)}(X,Y)) - D(p(X,Y) \| q^{(t+2)}(X,Y)) \leq \varepsilon.$$

stop and return $C_X = C_X^{t+2}$ and $C_Y = C_Y^{t+2}$, else go to step 2.

图 2 非线性相关失效数据联合聚类算法

由于非线性相关的失效数据联合聚类分析很难在可接受的时间内获得最优的聚类结果, 因此当互信息熵差小于某一给定任意小整数时, 可近似认为联合聚类达到局部最优. 为确保算法能在有限事件内获得联合聚类结果, 还需要保证失效数据联合聚类算法能够在有限迭代次数后收敛.

引理 1 非线性相关失效数据联合聚类算法总能在有限次迭代后收敛并达到局部最优.

证明 KL 散度可通过下式进行分解, 即

$$D(p(X,Y) \| q(X,Y)) = \sum_{\hat{x}} \sum_{x: C_X^{(t)}(x) = \hat{x}} p(x) \cdot \sum_y p(y|x) \log \frac{p(y|x)}{q^{(t)}(y|\hat{x})} \geq \sum_x \sum_{x: C_X^{(t)}(x) = \hat{x}} p(x) \cdot$$

$$\sum_y p(y|x) \log \frac{p(y|x)}{q^{(t)}(y|\hat{x})} = \sum_{\hat{x}} \sum_{x: C_X^{(t+1)}(x) = \hat{x}} p(x) \sum_{\hat{y}} \sum_{y: C_Y^{(t+1)}(y) = \hat{y}} p(y|x) \log \frac{p(y|x)}{q^{(t)}(y|\hat{y})} + \sum_x \sum_{x: C_X^{(t+1)}(x) = \hat{x}} p(x) \cdot \sum_{\hat{y}} \sum_{y: C_Y^{(t+1)}(y) = \hat{y}} p(y|x) \log \frac{1}{q^{(t)}(y|\hat{x})}.$$

而由于

$$\sum_x \sum_{x: C_X^{(t+1)}(x) = \hat{x}} p(x) \sum_{\hat{y}} \sum_{y: C_Y^{(t+1)}(y) = \hat{y}} p(y|x) \log \frac{1}{q^{(t)}(y|\hat{x})} \geq \sum_x q^{(t+1)}(\hat{x}) \sum_y q^{(t+1)}(\hat{y}|\hat{x}) \log \frac{1}{q^{(t+1)}(\hat{y}|\hat{x})}.$$

因此有

$$D(p(X,Y) \| q(X,Y)) \geq \sum_x \sum_{x: C_X^{(t+1)}(x) = \hat{x}} p(x) \cdot \sum_{\hat{y}} \sum_{y: C_Y^{(t+1)}(y) = \hat{y}} p(y|x) \log \frac{p(y|x)}{q^{(t+1)}(\hat{y}|\hat{x})q^{(t+1)}(y|\hat{y})} = D(p^{t+1}(X,Y) \| q^{t+1}(X,Y)).$$

由于每次迭代后行簇与列簇的数目是有限的, 且任意小整数 ε 确定, 所以非线性相关的失效数据联合聚类算法总能在有限次迭代后收敛且达到局部最优, 证明完毕.

3 实验结果与分析

3.1 失效特征提取

实验采用了美国洛斯阿拉莫斯国家实验室 (Los Alamos National Laboratory, LANL) 22 台高性能计算系统从 1996 ~ 2005 年的运行记录作为失效数据来源^[5]. LANL 高性能计算系统采用独立内存接口以及多级处理器结构, 共包含 4 750 个节点以及 24 101 个处理器, LANL 高性能计算机系统在规模、体系结构等方面均呈现出多样性, 大部分工作负载为 3D 仿真和可视化计算. 图 3 分别显示了 LANL 高性能计算机系统从 2003 年 9 月 1 日 ~ 2005 年 8 月 31 日 CPU 失效、内存失效、应用程序失效和文件系统失效 4 种类型的时间分布. 由图 3 可知, 不同的失效类型具有不同的时间分布特征, 在联合聚类分析时必须根据不同类型失效的特征分布进行区分.

3.2 联合聚类性能测试

实验给出了失效数据联合聚类算法在 LANL 失效数据集上运行的收敛速度和聚类效果, 同时引入迭代双聚类算法 (Iterative Double Clustering, IDC)^[11] 以及失效数据聚类的 temporal 和 spatial 算法进行比较. 算法在 Window XP 平台上实现, 硬件环境为 Intel Celeron CPU 2.4 GHz, 主存 1 G.

实验采用 LANL 实验室高性能计算机 2004 年 6 月~2005 年 9 月之间的失效数据作为聚类算法性能测试的数据来源.

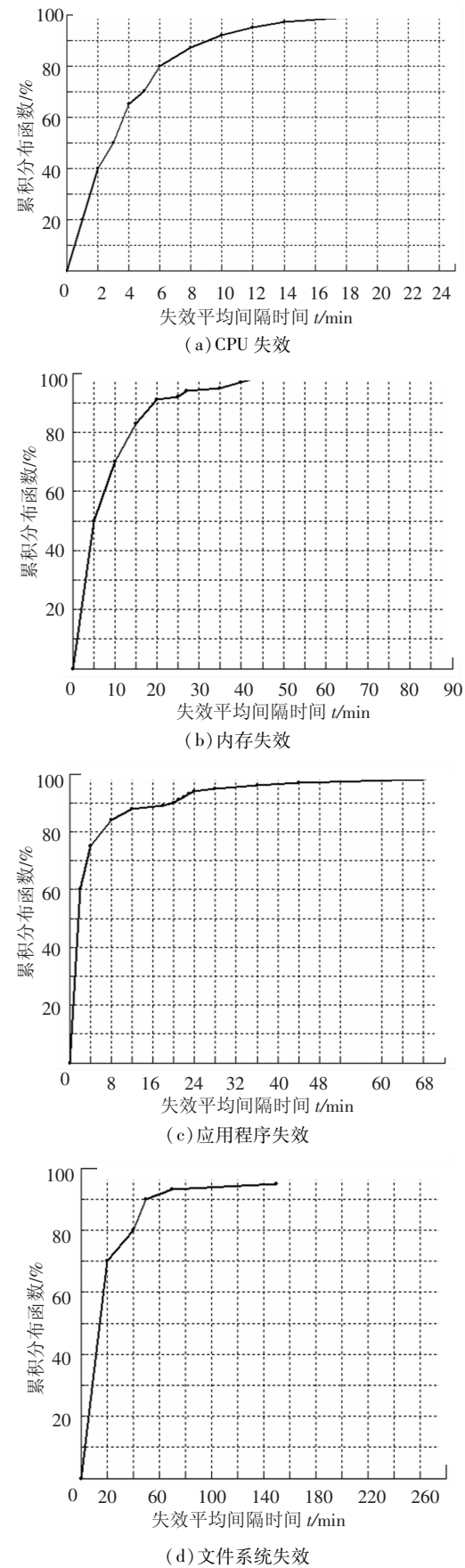


图3 4种失效类型时间分布

实验首先考察 4 类算法在不同数据规模下的运行时间. 图 4 为失效数据矩阵的列数固定时算法运行时间和数据行数的关系曲线, 分别测试矩阵行数 m 分别从 1 000 ~ 3 000 之间变化的运行时间. 图 5 为失效数据矩阵的行数固定时算法运行时间与矩阵维数的关系曲线, 分别测试维数从 5 ~ 25 的运行时间. 聚类分析时间随着失效样本数的增加而递增, 其中 temporal 和 spatial 算法耗时最短, 而迭代双聚类以及本文所提算法耗时较长, 这是因为前者属于单维聚类方法, 在聚类相似度计算时仅考虑行聚类问题, 而后者则从行聚类以及列聚类 2 方面进行聚类计算, 因此算法运行时间不可避免地延长.

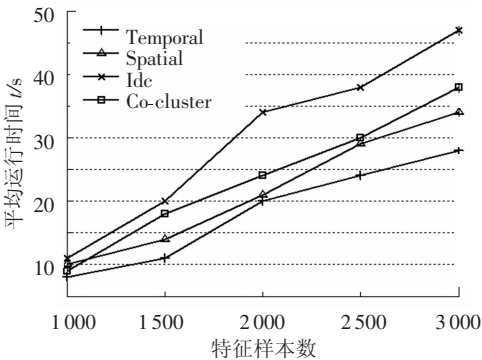


图4 列数固定时算法运行时间与行数关系曲线

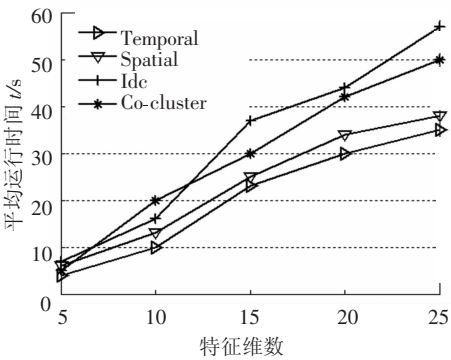


图5 行数固定时运行时间与列数关系曲线

图 6 表示在聚类过程中互信息熵差与算法运行时间的曲线关系. 由图 6 可知, 聚类过程中互信息熵差随着聚类算法运行时间递减, 在算法运行初始阶段, 2 类联合聚类算法熵差达到最大值, 在聚类分析 10 s 后开始平滑下降. 图 7 为聚类过程中算法迭代次数与运行时间的曲线关系. 由图 7 可知, 从聚类初始阶段开始, 当算法运行超过 30 s 后, 迭代也超过 50 次/s 且逐步达到最大值.

为评测聚类算法用于失效预测的有效性, 根据标注的 LANL 实验室失效数据真实类别计算聚类结果的查准率 (precision) 和查全率 (recall), 如表 1 所示. 表 1 的数据表明, 2 种联合聚类算法的聚类结果中 4 类失效类型的查准率均高于 65%. 这

是因为本文所提算法在单次行聚类或列聚类时均按照 KL 散度计算信息熵损失。作为对比,单维聚类算法由于仅按照时间分布或空间分布 2 种特征进行聚类,因此对于 4 种失效类型的聚类分析结果均不尽理想。对聚类过程中互信息熵差与迭代次数的考察说明,联合聚类算法能有效考虑不同特征的关联性,并采用关联特征来度量失效事件间相似性,因此联合聚类较单维聚类效果更好。

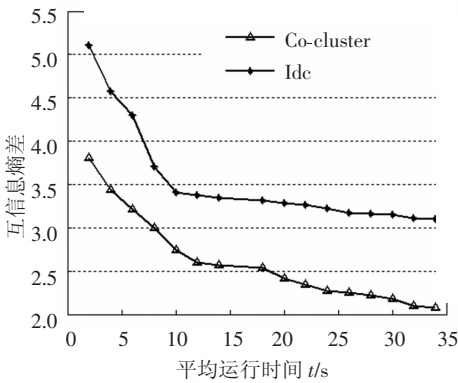


图 6 互信息熵差与运行时间的关系曲线

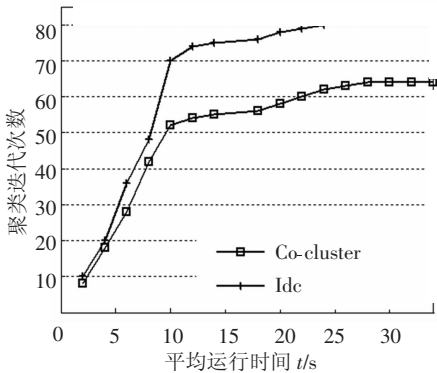


图 7 算法迭代次数与运行时间的关系曲线

表 1 聚类算法实验结果比较 %

	Co-cluster 算法		IDC 算法		Temporal 算法		Spatial 算法	
	查准率	查全率	查准率	查全率	查准率	查全率	查准率	查全率
CPU 失效	71.61	73.43	60.24	68.52	52.45	45.96	47.59	50.19
内存失效	65.72	70.24	62.48	65.45	43.81	54.64	54.46	55.89
应用程序失效	70.25	75.74	65.50	67.70	48.57	51.37	48.62	61.40
文件系统失效	68.44	77.10	66.00	70.83	45.82	55.92	50.05	59.83

4 结 论

1)针对高性能计算机系统非线性相关失效数据的高维、稀疏等特征,提出非线性相关失效数据联合聚类算法,以互信息熵损失差作为联合聚类度量标准并阐明了算法在有限次迭代后的收敛性。

2)实验计算了 4 种常见失效类型的时间分布,并比较了不同算法在失效数据集上的聚类效

果和收敛速度。实验结果表明,本文所提出的非线性相关失效数据联合在聚类准确性、聚类计算时间耗费等方面优于单向聚类方法。

参考文献:

[1] KEPHART J O, CHESS D M. The vision of autonomic computing [J]. IEEE Journal of Computer, 2003, 36(1): 41-50.

[2] SOLANO-QUINDE L D, BODE B M. Module prototype for online failure prediction for the IBM BlueGene/L [C]//IEEE International Conference on Elector/Information Technology. Ames: IEEE Computer and Communication society, 2008: 470-474.

[3] FU Song, XU Chengzhong. Exploring event correlation for failure prediction in coalitions of clusters[C]//Proceedings of the 2007 ACM/IEEE Conference on Supercomputing. New York, NY: ACM, 2007: 456-468.

[4] OLINER A, STEARLEY J. What supercomputers say: A study of five system logs[C]//Proceedings of the 37th Annual IEEE/IFIP International Conference on Dependable Systems and Networks. Washington, DC: IEEE Computer Society, 2007: 575-584.

[5] SCHROEDER B, GIBSON G A. A large-scale study of failures in high-performance computing systems [C]//Proceedings of the International Conference on Dependable Systems and Networks. Washington, DC: IEEE Computer Society, 2006: 249-258.

[6] LIANG Y L, ZHANG Y Y, SIVASUBRAMANIAM A A, et al. BlueGene/L failure analysis and prediction models [C]//Proceedings of the International Conference on Dependable Systems and Networks. Washington, DC: IEEE Computer Society, 2006: 425-434.

[7] LIANG Y L, SIVASUBRAMANIAM A, MOREIRA J. Filtering failure logs for a BlueGene/L prototype [C]//Proceedings of the 2005 International Conference on Dependable Systems and Networks. Washington, DC: IEEE Computer Society, 2005: 476-485.

[8] DHILLON I S, GUAN Y Q. Information theoretic clustering of sparse cooccurrence data [C]//Proceedings of the Third IEEE International Conference on Data Mining. Washington, DC: IEEE Computer Society, 2003: 517-520.

[9] JOHN G H, KOHAVI R, PFLEGER K. Irrelevant feature and the subset selection problem [C]//Proceedings of the 11th International Conference on Machine Learning. San Francisco: Morgan Kaufmann, 1994: 121-129.

[10] 闫雷鸣, 孙志挥, 吴英杰, 等. 联合聚类非线性相关的基因表达数据[J]. 计算机研究与发展, 2008, 45(11): 1865-1873.

[11] EL-YANIV R, SOUROUJON O. Iterative double clustering for unsupervised and semi-supervised learning [C]//Proceedings of the 12th European Conference on Machine Learning. London, UK: Springer-Verlag, 2001: 121-132.

(编辑 张 红)