

DOI: 10.7652/xjtub201703018

一种基于仿射传播的增强型流聚类算法

赵建龙¹, 曲桦^{1,2}, 赵季红², 蒋丁潮²

(1. 西安交通大学软件学院, 710049, 西安; 2. 西安交通大学电子与信息工程学院, 710049, 西安)

摘要: 针对目前流聚类算法无法有效处理数据流离群点的检测和处理, 以及增量式数据流聚类效率较低等问题, 提出了一种基于密度度量的异常检测、删除的增强型仿射传播流聚类算法。在仿射传播流聚类算法的基础上, 所提算法通过引进异常检测和删除机制改善了异常点对聚类精度、聚类效率的影响。利用仿射传播聚类实现在线数据流的聚类过程, 同时检测数据漂移现象, 即数据流分布特征随时间发生变化, 并采用基于密度度量的局部异常因子检测技术(LOF)对储备池数据进行异常检测和删除处理, 通过对当前类簇和处理过的储备池数据重聚类来重建动态数据流模型。在真实网络数据(KDD'99)上进行了实验, 结果表明, 所提算法不仅减少了重聚类构建动态模型的次数, 改善了聚类效率, 而且在同时考虑聚类精度、纯度和熵3种聚类评价标准下, 均优于传统的仿射传播流聚类算法。

关键词: 流聚类; 仿射传播; 局部异常因子; 异常删除

中图分类号: TP181 **文献标志码:** A **文章编号:** 0253-987X(2017)03-0105-06

An Enhanced Stream Clustering Algorithm Based on Affinity Propagation

ZHAO Jianlong¹, QU Hua^{1,2}, ZHAO Jihong², JIANG Dingchao²

(1. School of Software Engineering, Xi'an Jiaotong University, Xi'an 710049, China;

2. School of Electronics and Information, Xi'an Jiaotong University, Xi'an 710049, China)

Abstract: Aiming at the problem that the traditional stream clustering algorithm cannot effectively deal with the inspection and treatment of outliers, and the incremental data stream clustering efficiency is low, an enhanced stream clustering algorithm based on affinity propagation using density measurement was proposed. Based on the STRAP, the proposed algorithm can improve the clustering accuracy and efficiency by introducing a mechanism for outlier detection and removal. Firstly, the online stream clustering process is realized by the affinity propagation algorithm. Meanwhile, the phenomenon of data drift is detected, i. e., the distribution of data stream changes with time. In view of this phenomenon, the new algorithm can implement the outlier detection and removal in the reservoir based on local outlier factor, and then re-cluster the current cluster and the treated reservoir to reconstruct the dynamic stream clustering model. Finally, through the validation on the KDD'99 data, the experimental results showed that the proposed algorithm not only reduces the number of re-clustering and improves the clustering efficiency, but also is superior to the STRAP in terms of the three clustering evaluation criteria, i. e., the clustering accuracy, purity and entropy.

Keywords: stream clustering; affinity propagation; local outlier factor; outlier removal

收稿日期: 2016-05-28。 作者简介: 赵建龙(1992—),男,博士生;曲桦(通信作者),男,教授,博士生导师。 基金项目: 国家自然科学基金资助项目(61371087,61531013);国家“863计划”资助项目(2015AA015702)。

网络出版时间: 2016-12-27

网络出版地址: <http://www.cnki.net/kcms/detail/61.1069.T.20161227.1022.008.html>

数据流聚类需要对动态变化的数据对象实施快速增量处理并提供实时结果,不同于传统聚类可以多次访问静态数据,流聚类算法要求对数据流在内存和时间限制下实现单遍聚类过程。截至目前,聚类算法在数据流中的应用主要包括基于对象的流聚类和基于属性的流聚类^[1],其中基于对象的流聚类划分为在线数据抽象和离线聚类两个阶段^[2]。在线抽象阶段借助特定的数据结构来概述持续到达的数据流,同时通过定义时间窗口来减弱过早到达的数据流对于聚类结果的影响,常见的窗口模型包括滑动窗口模型^[2]、界标模型^[3]和阻尼模型^[4]。对于离线聚类阶段,通常采用传统的DBSCAN等聚类算法作为流聚类的基准算法实现离线聚类过程。基于属性的流聚类算法代表有ODAC、DGClust等在线分类聚类^[5-6]。这类算法的共性是对原始流量属性矩阵进行相关矩阵运算和聚类操作,由于现实中流量数据往往具有高维特征且随时间在不断变化,因此这类算法的核心任务转换为寻找一组不随时间而改变的最大相关近似属性集,从而简化算法的复杂度。

仿射传播算法(AP)^[7]作为一种基准聚类算法,很好地解决了K均值算法必须人工设定初始聚类参数的局限。文献[8]针对K均值聚类随机初始中心导致的结果不稳定问题,通过对聚类结果的加权集成提出了灵活高效的基于仿射传播的聚类集成算法。Zhang等突破性地将AP扩展到数据流挖掘领域,提出了在线流聚类算法(STRAP)^[9-10],通过增量式处理数据流,实现了在任意时刻进行聚类的目的。对比DenStream^[4],虽然STRAP聚类纯度和效率提高不少,但是在对持续快速到达的数据流进行聚类时,按照STRAP流程,算法会频繁地对当前类簇和储备池进行重聚类,从而导致聚类效率下降。另外,STRAP在初步判断异常点并将其存储到储备池过程中,选用基于距离的异常度量方式,而这种度量方式无法有效解决具有不同密度特征的数据集异常判断问题,即局部密度问题^[11],而且储备池中的数据仅仅是用于二次聚类,没有进一步认证和处理极限异常数据,这也在一定程度上影响了聚类效率。

本文通过增加对储备池异常数据的检测和删除过程,提出了增强型仿射传播流聚类算法(ESAP)。该算法采用一种简单高效的基于密度的局部异常因子(LOF)技术^[12]来检测储备池中的异常数据,然后通过删除具有较高局部异常因子的数据结点来降低储备池异常数据对聚类结果的影响,并实现降低重聚类次数和改善聚类效率的目的。

1 基于流聚类的仿射传播算法

AP是一种高效的基准聚类算法^[7],核心思想是以数据结点偏离类簇中心误差最小化为目标条件来寻找一组类代表点集合(聚类中心或exemplar)。给定数据集 $D = \{e_1, e_2, \dots, e_N\}$,其中 $e_i = (e_{i1}, e_{i2}, \dots, e_{in}) (i=1, 2, \dots, N)$ 表示 n 维空间中的一个数据点。AP算法将数据点对之间的相似度作为基准输入,相似度矩阵 $S = [s(i, k)]_{N \times N}$ 定义为

$$s(i, k) = \begin{cases} -d(i, k), & i \neq k \\ \sigma, & i = k \end{cases} \quad (1)$$

式中: $d(i, k)$ 为数据点 i 和 k 之间的欧式距离; σ 为数据点 i 成为聚类代表点的参考度, σ 越大表示数据点 i 成为聚类代表点的可能性也越大。

AP算法起始阶段将所有数据点当做潜在的聚类中心,通过吸引度 $r(i, k)$ 和归属感 $a(i, k)$ 两种信息传递不断收集点 i 成为合适聚类代表点的证据,最终通过迭代循环产生高质量的聚类中心簇 $C = \{c_1, c_2, \dots, c_N\}$,其中数据点 i 的聚类中心表示为

$$c_i = \arg \max_k \{r(i, k) + a(i, k), k = 1, \dots, N\} \quad (2)$$

AP最大优点在于无需人工设定聚类簇数,通过消息迭代循环来寻找最优类簇,然而直接将其应用到数据流挖掘中存在挑战。AP是基于批处理数据的挖掘算法,不具备在线增量数据挖掘能力,并且AP算法的时间复杂度为 $O(N^2)$,无法直接应用并高效处理大规模数据流聚类挖掘。

针对上述问题,STRAP利用加权聚类和分而治之的思想将AP应用到数据流挖掘中。加权近邻传播算法(WAP)^[9]主要目标是为了解决重复样本点的问题,区别在于重新定义相似度矩阵 $S' = [s'(i, k)]_{N \times N}$

$$s'(i, k) = \begin{cases} n_i s(i, k), & i \neq k \\ \sigma + (n_i - 1)\epsilon_i, & \epsilon_i \geq 0 \end{cases} \quad (3)$$

式中: n_i 为以数据点 i 作为其类代表点的数据点数; ϵ_i 为重复样本点间的平均距离。利用分而治之的思想提出分层加权近邻传播算法(Hi-WAP)^[10],将大规模数据集划分成多个小规模数据子集,分别利用AP进行聚类,然后对每个子集聚类所得的临时类簇进行WAP聚类得到流聚类模型,算法的复杂度从 $O(N^2)$ 降低到 $O(N^{1+\alpha})$, $\alpha > 0$,最后通过迭代更新聚类结果实现在线流聚类。

2 异常检测问题定义

数据流聚类需要能够快速检测异常值的存在,并能够采取恰当的措施。本文将针对储备池数据进行异常检测和删除操作,给出异常检测问题的相关定义。基于距离的异常检测方式根据数据点与近邻点间的距离远近程度来评定数结点的异常程度,这种度量方式简单高效,但却无法有效解决局部密度问题。本文探索性地采用基于密度的局部异常因子(LOF)来检测数据流中的离群点,LOF可以处理具有不同密度特征的数据集,使用数据点的局部偏离密度来表征其异常程度。给定储备池数据集 R ,分析针对 R 中数据点的 LOF 计算过程。

定义 1 对象 p 的 k 近邻距离。对于任意正整数 k ,对象 p 的 k 近邻距离表示为 $d_k(p)$,定义为对象 p 和 o 之间的距离 $d(p,o)$,必须同时满足以下两个条件:①数据集 R 中至少存在 k 个对象 $o' \in R \setminus \{p\}$ 满足 $d(p,o') \leq d(p,o)$;②数据集 R 中至多存在 $k-1$ 个对象 $o' \in R \setminus \{p\}$ 满足 $d(p,o') < d(p,o)$ 。其中 $\{p,o\} \in R$,距离度量方式可以是欧几里得距离、曼哈顿距离或其他任意一种相异性度量方式。

定义 2 $N_k(p)$ 为对象 p 的 k 近邻域,表示为

$$N_k(p) = \{q \mid d(p,q) \leq d_k(p), q \neq p\} \quad (4)$$

式中: $d(p,q)$ 为数据集 D 中对象 p,q 的直达距离。

LOF 是一种简单高效、基于密度的异常检测方法^[12],利用异常点周围的密度往往低于其近邻点密度的原理,即对象 p 的 $L(p)$ 越高表示其成为异常点的可能性越大。给定数据点 p 和正整数 k ,对象 p 的异常因子 $L(p)$ 表示为

$$L(p) = \frac{1}{|N_k(p)|} \sum_{q \in N_k(p)} \frac{D_k(q)}{D_k(p)} \quad (5)$$

式中: $d_k(p)$ 为对象 p 的直接可达密度,通过计算 $p,q \in N_k(p)$ 之间平均可达距离的倒数得到,即

$$D_k(p) = 1 / \left(\frac{\sum_{q \in N_k(p)} d_k(p,q)}{|N_k(p)|} \right) \quad (6)$$

一般来说,簇内对象的 L 约等于 1,簇边缘对象的 L 略大于 1,而离簇距离越远,对象的 L 就越大,也就是说该对象成为异常点的可能性更大。

3 基于仿射传播的增强型流聚类算法

3.1 算法描述

ESAP 算法在仿射传播流聚类基础上通过引进储备池数据的异常检测和删除机制实现了高效的数

据流聚类,ESAP 算法实现步骤如下。

步骤 1 利用 AP 聚类算法对初始数据流 $X = \{X_1, \dots, X_T\}$ 进行聚类得到初始化流聚类模型 CM,具体包括聚类簇数 K 和簇中心 $E_i (i=1, \dots, K)$ 。

步骤 2 对于持续到达的数据流 $X_t (t > T)$,计算其与当前各类簇中心间的距离 $d_{ij} (j=1, \dots, K)$,如果 $d_{ij} \leq \epsilon'$,将其合并到最近类簇并更新当前流模型。

步骤 3 如果 $d_{ij} \leq \epsilon'$,表明 X_t 不属于当前流模型,可能代表新模式或异常点,将其暂存至储备池 R 中等待进一步处理。

步骤 4 在处理持续到达的数据流同时检测数据漂移现象,即数据分布特征随时间是否发生变化。检测标准^[10]包括 PH 测试和限制储备池最大容量 M 。

步骤 5 如果满足 PH 测试或 $n_R \geq M$,则重启聚类过程用于构建适应新数据分布特征的流模型。重启聚类之前对储备池 $R = \{X_u, u=1, \dots, n_R\}$ 进行异常检测,计算每个结点的局部异常因子 $L_i (i=1, \dots, n_R)$,并按照 L 高低进行排序,然后删除具有较高局部异常因子的数据点得到更新的储备池 R' 。

步骤 6 通过对 R' 和当前类簇利用 AP 进行重聚类得到实时聚类模型,同时储备池 R 清空。

3.2 算法复杂度分析

ESAP 算法是建立在 STRAP 基础之上的,STRAP 利用分层加权聚类思想将 AP 的时间复杂度由 $O(N^2)$ 降低到 $O(N^{3/2})$,其中 N 代表数据流结点数。ESAP 增加了对储备池 R 的异常检测过程,而对检测方法 LOF 本身而言,单次运行时间复杂度主要取决于对 k 近邻数据库的检索。文献^[12]分析得出针对不同数据特性可采用不同的检索策略来简化时间复杂度。对于中维或中高维数据,采用索引结构使得 k -NN 的查询时间变为 $O(\log N)$,算法整体的时间复杂度为 $O(N \log N)$ 。本文的实验数据属于中高维数据,因此对大小为 M 的储备池进行单次 LOF 异常检测的时间复杂度为 $O(M \log M)$,假设在整个数据流聚类过程中重启聚类次数为 Δ ,那么执行 ESAP 的总时间复杂度为 $O(N^{3/2}) + \Delta O(M \log M)$,算法如下。

输入 初始流数据集 X ,距离阈值 ϵ'

输出 流聚类模型 CM

```
1 function ESAP( $X, \epsilon'$ )
2 CM ← AP( $X$ )
3  $R = \{\}$ 
4 for  $t > T$  do
5 计算  $X_t$  的最近类代表点  $E_i$ 
```

```

6      if  $d(x_i, E_i) < \epsilon'$  then
7          更新 CM
8      else
9           $R \leftarrow X_i$ 
10     end if
11     if 重聚类条件触发 then
12          $R \leftarrow \text{RL}(R, \epsilon, k)$ 
13          $\text{CM} \leftarrow \text{AP}(\text{CM}, R)$ 
14          $R = \{\}$ 
15     end if
16 end for
17 return CM
18 end function
19 function  $\text{RL}(R, \epsilon, k)$ 
20     for  $i = 1$  to  $n_R$  do
21         计算  $R[i]$  的  $k$  近邻  $N_k(R[i])$ 
22     end for
23     for  $i = 1$  to  $n_R$  do
24         计算  $R[i]$  的可直达密度  $d_k(R[i])$ 
25     end for
26     for  $i = 1$  to  $n_R$  do
27         计算  $R[i]$  的局部异常因子  $L(R[i])$ 
28     end for
29      $R_s \leftarrow \text{Sort}(R, L)$ 
30      $n_o \leftarrow \epsilon n_R$ 
31      $R' \leftarrow R_s[n_o + 1, \dots, n_R]$ 
32     return  $R'$ 
33 end function

```

4 实验分析

4.1 实验数据

本实验所采用的数据来源于网络入侵检测数据集 KDD'99^[10], 该数据集的主要应用目标是区别正常 TCP 连接和攻击, 每条记录包含一个局域网内某一时段的连接日志记录, 主要包括源 IP 地址、目的 IP 地址、持续时间、传输字节数等属性, 从源数据中抽取 34 维数值属性并进行处理和归一化, 得到 494 021 条流数据记录。每条连接对应的类标签包括正常类型(Normal)、拒绝服务攻击(DoS)、端口监视或扫描(Probe)、来自远程主机的未授权访问(R2L)和未授权的本地超级用户特权访问(U2R)5 类。

4.2 评价指标

在样本标签类的支持下, 本文采用的流聚类模

型评价指标包括准确度、纯度和熵^[13]。

准确度是一种简单、直观的聚类评价策略, 表示被准确分类样本数比例。评价本实验的聚类准确度时, 聚簇到同一类的样本点的类标签由其代表点的类标签决定, 即

$$A = 100 \frac{m}{n} \quad (7)$$

式中: m 为被正确聚类到流模型中的数据样本量; n 为所有样本数目。准确度反映了数据流聚类过程中被视作正常数据结点个数占所有数据量的比例, 准确度越高表示发现的离群点越少, 聚类模型更符合数据流分布特征。

聚类准确度的评价准则无法有效评价不平衡数据集的聚类问题, 并且对于小簇类的样本分类问题也表现平庸。纯度统计平均了所有簇中属于最大类的样本数目, 从而避免了采用准确度来评价不平衡数据集的聚类问题, 即

$$P = 100 \left(\sum_{i=1}^K \frac{|C_i^d|}{|C_i|} \right) / K \quad (8)$$

式中: K 为聚类数; $|C_i|$ 为类簇 i 的样本大小; $|C_i^d|$ 为在类簇 i 中最大类的样本数目。纯度越高表示被正确聚类的样本越多, 即聚类性能越佳。

熵是一种高效的聚类性能评价指标, 对于一个聚类 i , 首先需要计算每个类簇中样本的类分布, $p_{ij} = m_{ij} / m_i$ 表示聚类 i 中的成员属于类 j 的概率, m_i 表示聚类 i 中的所有成员个数, m_{ij} 表示聚类 i 中的成员属于类 j 的个数, 每个聚类的熵可表示为

$$E_j = - \sum_i p_{ij} \log(p_{ij}) \quad (9)$$

整体聚类的熵可表示为

$$E = \sum_{j=1}^m \left(\frac{N_j}{N} E_j \right) \quad (10)$$

式中: N 为数据集中总的样本数; N_j 为聚类 j 中的样本数。熵越低表示聚类效果越佳。

4.3 实验结果

本文提出的 ESAP 算法旨在通过增加异常检测和极限异常点的删除来强化在线流聚类算法, 实验思路包括两方面: 对比分析 AP 和 K 均值算法作为基准聚类算法在流数据挖掘中的性能; 对比分析 ESAP 和 STRAP 实现流聚类挖掘中的聚类效率和聚类质量。

流聚类准确度对比结果如图 1 所示, 其中 STRKC 是在 STRAP 基础上将 AP 基准聚类算法替换成 K 均值而实现的在线流聚类算法。ESAP(5) 和 ESAP(10) 分别表示储备池极限异常删除比例为

5%和 10%的 ESAP 算法,初始化流模型的参数设置包括初始时间序列点 $T=1\ 000$,储备池最大容量 $M=300$ 。由图 1 可知,随着时间序列的延长,4 条曲线均呈下降趋势,但无论是 STRAP 还是 ESAP,整体的聚类准确度均高于 STRKC,说明 AP 作为基准聚类算法优于 K 均值算法。

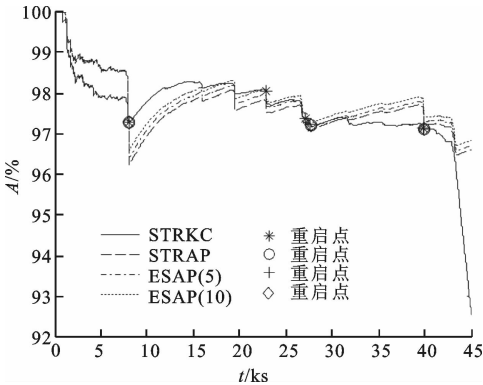


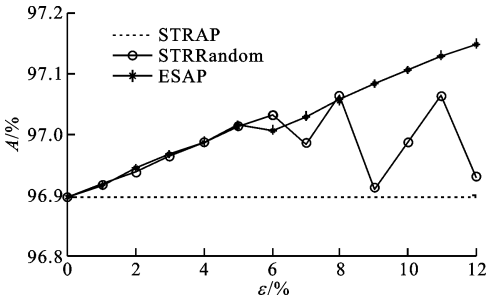
图 1 流聚类准确度对比结果

基于聚类准确度、纯度和熵 3 种评价指标的聚类结果对比如图 2 所示,实验出发点是考虑删除比例对于 STRAP 和 ESAP 聚类质量的影响。另外,考虑到删除数据结点是否能对聚类结果产生影响,实验增加了随机删除储备池数据结点并进行了聚类质量检验,表示为 STRRandom。由图 2 可知,随机删除储备池数据结点,准确度、纯度和熵值曲线变化不大,流聚类性能不受随机数据点删除的影响;随着储备池极限异常删除比例的提升,聚类准确度、纯度呈现增长趋势,而且熵值曲线也呈下降趋势,表示聚类性能随着删除比例的增加,表现更加良好,从而可知 ESAP 比 STRAP 具有更优的聚类性能。

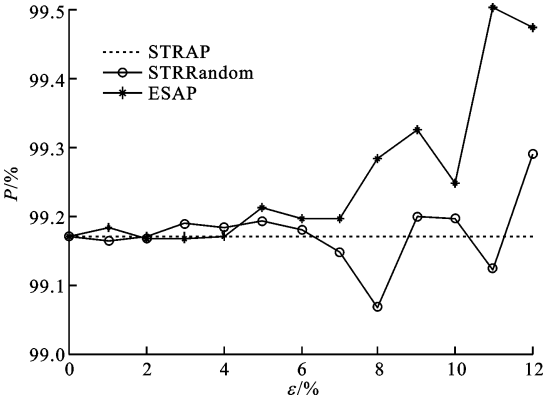
图 1 通过对比任意时刻的流聚类准确度结果,ESAP 的准确度整体高于 STRAP,表明 ESAP 更能反映出当前流数据聚类模型。另外,ESAP(10)曲线整体也优于 ESAP(5),表明删除比例越高聚类结果越佳。图 2 在给定时间序列长度的情况下展现了储备池异常删除比例对整体聚类性能的影响,反映出随机删除储备池数据点对整体聚类结果没有太大影响,而 ESAP 随着删除比例的增加,聚类性能呈增长趋势,表明聚类准确度、纯度越高,熵值越低。

5 结 论

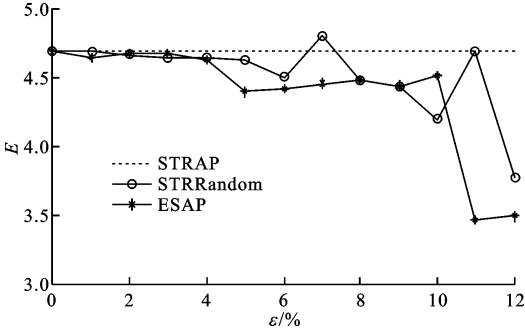
本文提出了一种基于密度度量的异常检测和删除的增强型仿射传播流聚类算法,通过增加储备池异常检测和极限异常结点的删除过程来增强流聚类的效率,并改善聚类质量。通过真实网络数据流聚



(a) 聚类准确度对比



(b) 聚类纯度对比



(c) 聚类熵值对比

图 2 聚类结果对比图

类验证可知,在同时考虑聚类准确度、纯度和熵值等评价标准时,ESAP 的聚类效果优于传统仿射传播流聚类。未来的研究工作将继续探索使用其他类型的异常检测方法,并寻找一种决定最优极限异常点删除比例的标准来优化流聚类模型。

参考文献:

[1] SILVA J A, FARIA E R, BARROS R C, et al. Data stream clustering: a survey [J]. ACM Computing Surveys, 2013, 46(1): 125-134.
[2] ZHOU Aoying, CAO Feng, QIAN Weining, et al. Tracking clusters in evolving data streams over sliding windows [J]. Knowledge and Information Systems,

- 2008, 15(2): 181-214.
- [3] ACKERMANN M R, MARTENS M, RAUPACH C, et al. StreamKM++: a clustering algorithm for data streams [J]. *Journal of Experimental Algorithmics*, 2012, 17: 173-187.
- [4] CAO Feng, ESTER M, QIAN Weining, et al. Density-based clustering over an evolving data stream with noise [C] // *SIAM International Conference on Data Mining*. Philadelphia, PA, USA: SIAM, 2006: 328-339.
- [5] RODRIGUES P P, GAMA J, PEDROSO J P. ODAC: hierarchical clustering of time series data streams [C] // *Proceedings of the 2006 SIAM International Conference on Data Mining*. Philadelphia, PA, USA: SIAM, 2006: 615-627.
- [6] GAMA J, RODRIGUES P P, LOPES L M B. Clustering distributed sensor data streams using local processing and reduced communication [J]. *Intelligent Data Analysis*, 2011, 15(1): 3-28.
- [7] FREY B J, DUECK D. Clustering by passing messages between data points [J]. *Science*, 2007, 315 (5814): 972-976.
- [8] 王羨慧, 覃征, 张选平, 等. 采用仿射传播的聚类集成算法 [J]. *西安交通大学学报*, 2011, 45(8): 1-6.
- WANG Xianhui, QIN Zheng, ZHANG Xuanping, et al. Cluster ensemble algorithm using affinity propagation [J]. *Journal of Xi'an Jiaotong University*, 2011, 45(8): 1-6.
- [9] ZHANG Xiangliang, FURTELEHNER C, SEBAG M. Data streaming with affinity propagation [C] // *Lecture Notes in Computer Science*. Berlin, Germany: Springer, 2008: 628-643.
- [10] ZHANG Xiangliang, FURTELEHNER C, GERMAIN-RENAUD C, et al. Data stream clustering with affinity propagation [J]. *IEEE Transactions on Knowledge & Data Engineering*, 2014, 26(7): 1644-1656.
- [11] HA J, SEOK S, LEE J S. Robust outlier detection using the instability factor [J]. *Knowledge-Based Systems*, 2014, 63(2): 15-23.
- [12] BREUNIG M M, KRIEGEL H P, NG R T, et al. LOF: identifying density-based local outliers [J]. *ACM Sigmod Record*, 2000, 29(2): 93-104.
- [13] WU Junjie, HUI Xiong, CHEN Jian. Adapting the right measures for K-means clustering [C] // *Proceedings of the 15th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. New York, USA: ACM, 2009: 877-886.

(编辑 赵炜)

(上接第 86 页)

- [9] 王堃, 孙勇, 彭明军, 等. 基于 ANSYS 的铝蜂窝夹芯板低速冲击仿真模拟研究 [J]. *材料导报*, 2012, 26 (4): 157-160.
- WANG Kun, SUN Yong, PENG Mingjun, et al. Simulation of low-velocity Impact of aluminium honeycomb sandwich panels with ANSYS [J]. *Mater Rev*, 2012, 26(4): 157-160.
- [10] QIAO Jinxiu, CHEN Changqing. In-plane crushing of a hierarchical honeycomb [J]. *International Journal of Solids and Structures*, 2016, 85/86: 57-66.
- [11] ZHANG Yong, LU Minghao, WANG C H, et al. Out-of-plane crashworthiness of bio-inspired self-similar regular hierarchical honeycombs [J]. *Composite Structures*, 2016, 144: 1-13.
- [12] ASHAB A, RUAN D, LU G, et al. Quasi-static and dynamic experiments of aluminum honeycombs under combined compression-shear loading [J]. *Materials & Design*, 2016, 97: 183-194.
- [13] 张健, 赵桂平, 卢天健. 泡沫金属在冲击载荷下的能量吸收特性 [J]. *西安交通大学学报*, 2013, 47(11): 105-112.
- ZHANG Jian, ZHAO Guiping, LU Tianjian. Energy absorption behavior of metallic cellular materials under impact loading [J]. *Journal of Xi'an Jiaotong University*, 2013, 47(11): 105-112.
- [14] 李响, 周幼辉, 童冠, 等. 超轻多孔类蜂窝夹心结构创新构型及其力学性能 [J]. *西安交通大学学报*, 2014, 48(9): 88-94.
- LI Xiang, ZHOU Youhui, TONG Guan, et al. Innovating configuration and mechanical properties of the core for ultralight and porous quasi-honeycomb sandwich structure [J]. *Journal of Xi'an Jiaotong University*, 2014, 48(9): 88-94.

(编辑 葛赵青)