

基于深度文本聚类的关键词语义多模态挖掘方法

周春良¹ 杨畅畅¹ 李晓辉¹

ZHOU Chunliang YANG Changchang LI Xiaohui

摘要

多模态数据具有信息丰富性,但单一挖掘方法往往难以全面把握其互补、冗余和协同的关联,导致结果不精确或片面。针对这一问题,文章提出了一种新颖的基于深度文本聚类的关键词语义多模态挖掘方法。首先,利用深度学习的半监督学习算法,精准抽取并有效融合多模态数据中的语义特征。接着,通过结合 BWP 指标和 FCM 聚类法,确定最佳聚类数,并实现对这些特征的精准聚类分析,形成语义聚类簇。最后,借助语义映射技术,将不同模态的特征映射到同一语义空间,从而全面挖掘关键词的深层语义信息。实验结果表明,所提方法在关键词语义挖掘任务中表现优异,BWP 值最高达到了 0.9,R@1、R@5 和 R@10 指标分别高达 0.985、0.974 和 0.969,挖掘性能良好。

关键词

深度文本聚类;半监督学习算法;BWP 指标;语义映射;关键词语义挖掘

doi: 10.3969/j.issn.1672-9528.2025.03.037

0 引言

多模态数据挖掘涉及处理多种类型或来源的数据,其中包括文本、图像、视频、音频等不同形式和结构的数据。由于数据的多样性和复杂性,多模态数据挖掘易忽视关联性,导致部分数据丢失。

刘长红等人^[1]通过图卷积网络学习文本数据间的区域关系并采用双向门控循环单元建立文本单词间的关系;利用张量融合网络对这些语义关系展开匹配,学习多模态数据间的细粒度语义关联;在此基础上通过门控循环单元学习文本关键词的全局特征,实现语义多模态的挖掘。该方法需要大量的计算资源和时间来处理大规模数据集,实时性较差。Zheng 等人^[2]基于位置敏感哈希和 Bloom Filter 创建文档索引向量和查询向量,用该方法获取编辑距离作为多模态文本数据的关键词差异标准;在此基础上将关键词标准拆分为多个单词形式,对其分别展开 top-k 语义排序搜索,以此完成多模态文本数据之间的关键词语义挖掘。该方法主要关注文本模态,缺少对多模态数据语义理解的片面性。杨帆等人^[3]首先通过引入多模态双重注意力机制的文本语义增强模块对多模态文本数据语义特征实施语义增强;其次在该模块中加入保留强度和更新强度功

能,以此控制多模态语义特征中查询关键词语义特征的保留和更新程度;最后对其展开一定优化,已完成整体多模态文本数据的关键词语义挖掘。该方法采用的多模态双重注意力机制,在未见过的数据分布或任务上容易面临泛化能力不足的问题。Hussain 等人^[4]采用 MaxNet 模型从多模态文本数据中提取深层特征,并通过分层堆叠的方式持续更新模型。同时,实施聚类处理来整合各种特征,并引入 softmax 概率来计算聚类后特征的语义相似性以实现语义挖掘。当特征之间的差异显著时,这种方法所使用的 softmax 概率可能无法精确地度量它们之间的相似性。

因此,本文提出一种基于深度文本聚类的关键词语义多模态挖掘方法。首先,通过深度学习的半监督学习算法,精确捕捉并融合多模态数据中的关键语义特征。其次,利用 BWP 指标与 FCM 聚类法相结合,科学确定最佳聚类数,确保数据的精准分类。随后,采用语义映射技术,将不同模态的特征统一映射至同一语义空间,从而深入挖掘关键词的深层语义信息。此技术路径不仅充分发挥了多模态数据的互补优势,还提升了关键词语义挖掘的准确性和全面性。

1 深度文本聚类

1.1 基于深度学习的多模态语义特征抽取和融合

针对关键词语义多模态文本数据环境中的数据结构异构特性^[5],提出基于深度学习的半监督学习算法,从原始的多模态文本数据中抽取得到具有代表性的深度语义特征,并降维到低维空间,简化数据的复杂性。网络模型结构如图 1 所示。

1. 郑州西亚斯学院电信与智能制造学院 河南郑州 450000
[基金项目] 郑州西亚斯学院 2024 年度校级科研项目 (2024XKD077); 河南省本科高校 2023 年度产教融合研究项目; 郑州西亚斯学院应用型本科高校建设专项教改项目 (2022YYXZX01)

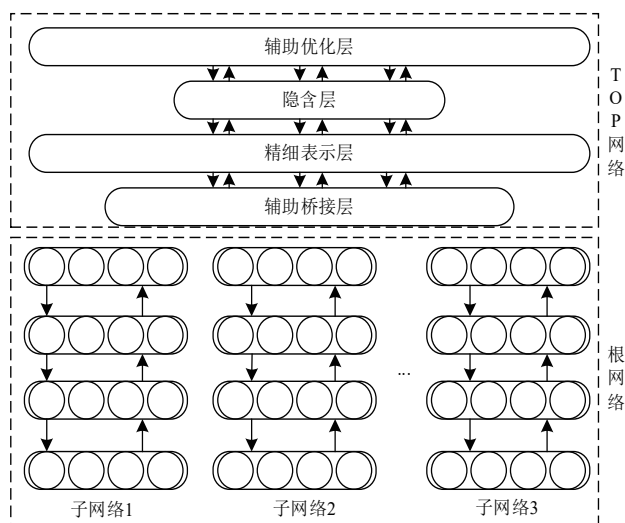


图1 半监督多模态神经网络总体结构图

该网络分为两个主要部分：根部网络与上层网络。

(1) 根部神经网络的同构特征抽取

将自动编码器^[6]作为网络的基本结构对根部网络展开训练，过程分为无监督预训练和有监督多模态联合微调阶段。

输入多模态文本数据 x ，在训练每一个隐藏层时通过权重矩阵 E_1 和映射函数 $\gamma(\cdot)$ 将 x 作非线性变换处理，并通过编码映射后，获取其隐含特征表达 $j = \gamma(E_1^T x)$ ，在解码阶段通过新的权重矩阵 e_2 和 $\gamma(\cdot)$ 将原始的输入 x 重构，获取重构后的输入 $x' = \gamma'(e_2^T j)$ ，将 x 与 x' 之间的重构误差作为优化目标展开无监督训练：

$$\text{Loss}(x, x') = \|x - x'\|_2^2 \quad (1)$$

完成独立模态的预训练后，对整个网络实施多模态联合微调。为了分析和保留各模态间的关系，在子网络最上层加入共享的辅助层，并与预训练的子网络一同调整参数。训练中令权值连接 Y_m 共享一个权值矩阵，以抽取相似的模态语义特征，提高多模态语义特征抽取的效率和准确性，这个过程采用引入最小化重构误差损失函数的反向传播算法^[7]实现，公式为：

$$Y = -\sum_j \sum_i \lg \left(A \left(U = u^{(i)} \middle| j_m^{(i)}, Y, b_{\text{root}} \right) \right) \quad (2)$$

式中： j_m 表示第 m 个模态的最上层神经元； N 表示训练样本 $x^{(i)}$ 的数量； $u^{(i)}$ 表示 $x^{(i)}$ 的类标； $j_m^{(i)}$ 表示第 m 个子网络对于输入 $x^{(i)}$ 的顶层输出； Y 表示权值矩阵； b_{root} 表示偏置向量。在此基础上引出 j_m 的 k 类分类问题，将 j_m 作为输入向量，求出 j_m 属于类别 i 的概率实现特征分类：

$$A \left(U = u^{(i)} \middle| j_m^{(i)}, Y, b_{\text{root}} \right) = \frac{\gamma \exp(Y^i j_m + b_{\text{root}}^i)}{\sum_{\varphi} \exp(Y^{\varphi} j_m + b_{\text{root}}^{\varphi})} \quad (3)$$

式中： Y^{φ} 表示权值矩阵 Y 的第 φ 个行向量。

最终采用梯度下降优化策略对整个模型参数展开相应调

整，以便从各子网络顶层隐藏层中提取高级抽象特征表达 j_m 。

(2) 上层网络半监督融合特征学习

在获取各个模态的高级抽象特征后，这些特征在上层网络中整合，以抽取最终的多模态共享融合特征。上层网络包括输入层、隐藏层和辅助优化层，根部网络获取的特征表达 j_m 将转换为上层网络的输入 x' 经过 $j = s(Ex' + b)$ 转换投影至低维特征空间。辅助优化层优化权值矩阵 Y 的参数，过程涉及 Logistic 分类器^[8]对顶部三层网络的训练，损失函数采用有监督判别损失函数 Y_{dis} 和无监督生成损失函数 Y_{gen} ， Y_{gen} 用于衡量重构后的输入 x' 与输出 $\hat{x}'$ 之间的重构损失，若值越小，则表明抽取到的融合特征中保留的原始信息越多，公式为：

$$\begin{cases} Y = Y_{\text{dis}} + Y_{\text{gen}} + \mu_1 \|E\|_{\varphi} + \mu_2 \|E\|^2 \\ Y_{\text{dis}} = -\sum_i \lg \left(A \left(U = u^{(i)} \middle| j^{(i)}, B, b_{\text{top}} \right) \right) \\ Y_{\text{gen}} = -\sum_i \left[x^{(i)} \lg \hat{x}^{(i)} + (1 - x^{(i)}) \lg (1 - x^{(i)}) \right] \end{cases} \quad (4)$$

式中： B 表示数据真实标签； μ_1 、 μ_2 表示原始特征维度和目标特征维度； b_{top} 表示偏置向量。最终得到 Y_{gen} 最小化的输出 $\hat{x}'$ ：

$$\hat{x}' = s(E^Y s(Ex' + b) + b') \quad (5)$$

式中： $s(\cdot)$ 表示 sigmoid 函数； b 和 b' 表示偏置项，这些特征融合了不同模态的深度语义信息。

1.2 多模态文本聚类

1.2.1 最佳聚类数确定

聚类的理想结果是类内紧密而类间远离，为此采用簇间-簇内百分比 (between-within percentage, BWP) 指标作为最佳聚类目标，该指标不受量纲影响，综合考量类内距离最小化和类间距离最大化。

令 w 为聚类空间，假设 n 个多模态文本深度语义特征 $\hat{x}'$ 被聚类为 l 类，则定义第 o 类的第 i 个语义特征样本 $\hat{x}'$ 的 BWP 指标为^[9]：

$$\begin{cases} L_{\text{BWP}}(\hat{x}_o', x_i') = (v(o, i) - q(o, i)) / (v(o, i) + q(o, i)) \\ v(o, i) = \min_{1 \leq a \leq l, a \neq o} \left(\frac{1}{n_a} \sum_{p=1}^{n_a} \|x_p^{(a)} - x_i^{(o)}\|^2 \right) \\ q(o, i) = \frac{1}{n_o - 1} \sum_{q=1, q \neq i}^{n_o} \|x_q^{(o)} - x_i^{(o)}\|^2 \end{cases} \quad (6)$$

式中： $v(o, i)$ 、 $q(o, i)$ 分别表示第 o 类的第 i 个语义特征样本的最小类间距离和类内距离； $\|\cdot\|^2$ 表示平方欧氏距离； $x_i^{(o)}$ 和 $x_p^{(a)}$ 分别表示第 o/a 类的第 i/p 个样本，其中 a 和 o 为类标； n_a 表示第 a 类中的样本个数； $x_q^{(o)}$ 表示第 o 类的第 q 个语义特征样本，且 $q \neq i$ ； n_o 表示第 o 类中的语义特征样本个数。

BWP 指标值越大表明单个语义特征样本的聚类效果越好, 为此将 K-means 算法与该指标相结合, 在给定聚类数搜索范围 $[l_{\min}, l_{\max}]$ 的基础上, 求出整个多模态深度文本语义特征中所有语义特征样本的 BWP 指标平均值 $L_{\text{avgBMP}}(l)$ 即可确定最佳聚类数 l_{opt} , 公式为:

$$\begin{cases} L_{\text{avgBMP}}(l) = \frac{1}{n} \sum_{o=1}^l \sum_{i=1}^{n_o} (\hat{x}_o^y, \hat{x}_i^y) \\ l_{\text{opt}} = \arg \max_{2 \leq l \leq n} \{L_{\text{avgBMP}}(l)\} \end{cases} \quad (7)$$

1.2.2 FCM 聚类算法

确定最佳聚类数 l_{opt} 后, 采用 FCM 聚类算法结合最佳聚类数进一步优化聚类结果, 该算法将 n 个多模态深度文本语义特征 $\hat{x}' = \{x_1, x_2, \dots, x_n\}$ 划分为 c 个模糊组, 求得模糊组 i 的聚类中心 c_i 以实现目标函数非相似性指标最小, 目标函数公式为:

$$\begin{cases} K(\mathbf{O}, c_1, c_2, \dots, c_c) = \sum_{i=1}^c K_i = \sum_{i=1}^c \sum_{o=1}^n o_{io}^y d_{iq}^2 \\ \mathbf{O} \in \sum_{i=1}^c o_{io} = 1, \forall i = 1, 2, \dots, n \end{cases} \quad (8)$$

式中: $d_{iq} = \|\mathbf{c}_i - \mathbf{x}_q\|$ 表示第 i 个聚类中心和第 q 个特征样本点之间的欧式距离; $\mathbf{O}$ 表示隶属度矩阵, o_{io} 表示第 i 个聚类中心第 o 类的隶属度; $y \in [1, \infty)$ 表示加权指数。

为了得到最佳模糊聚类划分, 在给定最佳聚类数的基础上, 求得 $\min\{K(\mathbf{O}, c_1, c_2, \dots, c_c)\}$, 约束条件为 $v_o(o = 1, 2, \dots, n)$, 新目标函数为:

$$\begin{aligned} \bar{K}(\mathbf{O}, c_1, c_2, \dots, c_c, v_1, v_2, \dots, v_n) = \\ K(\mathbf{O}, c_1, c_2, \dots, c_c) + \sum_{o=1}^n v_o \left(\sum_{i=1}^n o_{io} - 1 \right) \end{aligned} \quad (9)$$

最后, 根据上式求出 c 个聚类中心:

$$c_i = \left(\sum_{o=1}^n o_{io}^a x_i \right) / \sum_{o=1}^n o_{io}^a \quad (10)$$

若目标函数式 (9) 的结果小于某个确定的阈值则停止算法; 反之则采用下式求出新的 $\mathbf{O}$ 矩阵, 并通过式 (10) 持续迭代, 直到符合终止条件为止, 以此确定每个多模态深度文本语义特征的分类, 将相似的特征点聚类成一组, 形成多个聚类簇。

2 关键词语义多模态挖掘

针对每个聚类簇, 运用基于语义分析的映射方法, 通过挖掘多模态文本特征间的语义标签信息对同模态文本展开提取, 并揭示不同模态文本之间的语义关联。

(1) 多模态数据表达

引入基于多模态文本特性和多模态语料的数据表达机

制-多模态语义文档 (MAD), 统一表达不同模态的文本数据, 并反映其潜在语义关系, 命名相似语义的多模态文本为多模态语义文档。

由聚类簇中多类模态深度文本语义特征组成一个 L_{MAD} 集合, 通过 N_v 个语义类别表示为:

$$L_{\text{MADocs}} = \{L_{\text{MAD}1}, L_{\text{MAD}v}, L_{\text{MAD}N_v}\} \quad (11)$$

其中语义类别为 v 的 $L_{\text{MAD}v}$ 为:

$$L_{\text{MAD}v} = \bigcup_{\#} \{t_i^{\#}(v) | v, i = 1, 2, \dots, I_v^{\#}\} \quad (12)$$

式中: $t_i^{\#}(v)$ 表示从 v 中第 i 个多模态数据实例中提取的描述集合; $I_v^{\#}$ 表示 v 中包含的不同模态总数。

(2) 多模态语义融合与相关性计算

使用多模态高斯 PLSA 模型 (GM-PLSA) 和不对称算法建立同构的潜在主题语义空间, 能够有效地处理多模态数据之间的不平衡问题, 并深入挖掘多模态数据的关键词语义内涵, 从而更准确地挖掘出与关键词相关的语义信息。

在模型中, 多模态语义文档 $L_{\text{MAD}v}(v \in 1, \dots, N_v)$ 和多模态底层特征向量 $\mathbf{x}_i^{\#}$ 之间的关联性用潜在主题变量 $d_k(k \in 1, 2, \dots, K)$ 描述^[10], 即:

$$h(L_{\text{MAD}v}, \mathbf{x}_i^{\#}) = \sum_{d_k} h(d_k) h(L_{\text{MAD}v} | d_k) h(\mathbf{x}_i^{\#} | d_k) \quad (13)$$

d_k 与 $\mathbf{x}_i^{\#}$ 之间存在的条件概率为:

$$h(\mathbf{x}_i^{\#} | d_k) = 1 / \left(2\pi^{l_{\text{Dim}}/2} |\sum_k|^{1/2} \right) \exp \frac{(\mathbf{x}_i^{\#} - \mathbf{\kappa}_i^{\#})^T \sum_k^{-1} (\mathbf{x}_i^{\#} - \mathbf{\kappa}_i^{\#})}{2} \quad (14)$$

式中: l_{Dim} 表示 $\mathbf{x}_i^{\#}$ 的维数; $\sum_k$ 表示 $l_{\text{Dim}} \times l_{\text{Dim}}$ 协方差矩阵; $\mathbf{\kappa}_i^{\#}$ 表示 l_{Dim} 维均值向量。

运用不对称学习算法对各类模态的关键词语义信息实施融合。首先, 通过 GM-PLSA 模型对某一模态展开学习, 以实现对其他模态的互补和优化。其次采用期望最大化算法 (EM) 估算 $h(d_k)$ 、 $h(L_{\text{MAD}i} | d_k)$, 接着学习另一种模态, 并保持 $h(d_k)$ 、 $h(\mathbf{x}_i^{\#} | d_k)$ 不变; 其次在已知模态特征分布的条件下, 学习某主题条件下的特征向量分布 $h(\mathbf{x}_i^{\#} | d_k)$; 最后从语义上将多模态数据融合, 计算出该数据的主题分布 $h(\mathbf{x} | L_{\text{MADnew}})$ 。

基于前述模型训练所得的主题分布, 并运用余弦夹角公式来获取相关性结果, 用 $L_{\text{MAD}i}$ 、 $L_{\text{MAD}j}$ 分别表示两个不同模态的对象, 其相关性的计算公式为:

$$C(L_{\text{MAD}i}^{\#}, L_{\text{MAD}j}^{\#}) = \frac{\sum_{d_k} h(d_k | L_{\text{MAD}i}^{\#}) \cdot h(d_k | L_{\text{MAD}j}^{\#})}{\sum_{d_k} \left[h(d_k | L_{\text{MAD}i}^{\#}) \right] \cdot \left[h(d_k | L_{\text{MAD}j}^{\#}) \right]} \quad (15)$$

当 $L_{\text{MAD}i}$ 、 $L_{\text{MAD}j}$ 在语义上越相似或关联性越高时, 它们的主题分布向量之间的夹角余弦值越接近于 1, 这意味着它们的关键词语义模态越相似, 挖掘出多模态数据间的语义关联, 获得高度语义相关性的多模态信息。

3 实验与分析

为了验证基于深度文本聚类的关键词语义多模态挖掘方法的整体有效性,需要对其展开测试。实验平台的配置包括 Ubuntu 16.04.5 LTS 64 位操作系统,搭载 Intel(R)Core(TM)i7 处理器,主频为 3.50 GHz,拥有 2 核处理器,配备 8 GB 显存和 4.25 TB 的硬盘容量。

实验数据集从 Wikipedia 生成,包含 2 173 个训练文本和 693 个测试文本对,每种特征作为一个数据模态,包括文本模态(单词 a1、词性 a2、语法 a3)、图像模态(像素 b1、纹理 b2、物体检测 b3)、音频模态(波形 c1、梅尔频谱 c2、声谱图 c3)以及噪声模态(带噪音特征 n1、均匀分布噪声 n2、高斯分布噪声 n3)。观察使用所提方法得到的模态权值信息,结果如图 2 所示。

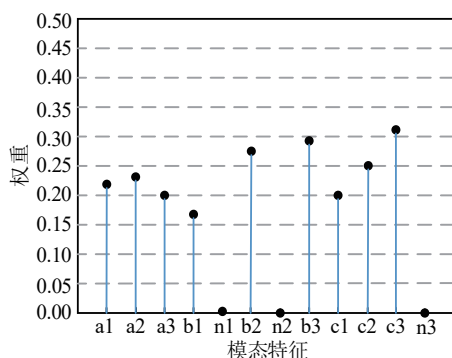


图2 不同方法下各模态特征权重

分析图 2 可知,经过所提方法处理后,噪声模态(带噪音特征 n1、均匀分布噪声 n2、高斯分布噪声 n3)均为 0,表明该方法能够有效过滤噪声,并为关键词语义挖掘任务相关的模态赋予较高权值,证明了该方法的有效性。

以 BWP 值作为评估指标,验证所提方法对整体数据集的聚类效果,值越大表明聚类结果越好,对比结果如图 3 所示。

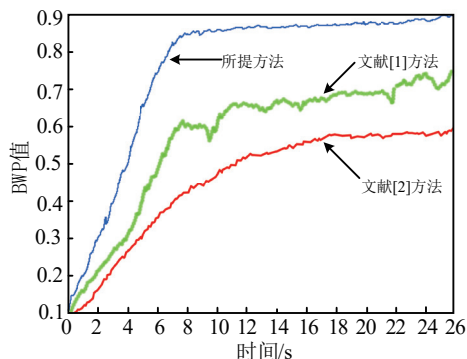


图3 所提方法聚类效果

分析图 3 可知,所提方法的 BWP 值随着时间的增加稳步上升,并在第 9 s 时趋于稳定,最高达到了 0.9。相比之下,其他方法的 BWP 值整体上升趋势较慢且波动较大,文献 [1]

方法最高为 0.75,文献 [2] 方法最高为 0.6。这说明所提方法具有较快的聚类效率,能够在短时间内准确稳定地实现多模态文本数据特征的聚类。

使用 R@1 (第一名准确率)、R@5 (前五名准确率)和 R@10 (前十名准确率)作为评价指标,评估所提方法的挖掘性能。R@1、R@5 和 R@10 是用于衡量挖掘出的关键词与原始文本之间的语义相似度的关键指标,分别表示挖掘出的关键词语义中排名最高、前五、前十分别与原始文本语义最相关的关键词的占比,这些指标越高,说明挖掘出的关键词与原始文本的语义相似度越高,因此可以认为该方法的性能越好,计算公式为:

$$\begin{cases} R@1 = \frac{L_1}{W} \\ R@5 = \frac{L_5}{W} \\ R@10 = \frac{L_{10}}{W} \end{cases} \quad (16)$$

式中: L_1 、 L_5 、 L_{10} 分别表示挖掘出的关键词列表中排名第一、前五、前十的关键词与原始文本语义相关的次数; W 表示总测试次数。

对比结果如表 1 所示,由表中的数据可以看出,所提方法在 R@1 上达到了 0.985 的高分,相较于文献 [1] 和文献 [2] 方法分别提高了约 14 和 20 个百分点,在 R@5 和 R@10 上,所提方法同样展现出较高的性能,分别达到了 0.974 和 0.969,表明所提方法在捕捉最相关关键词方面表现卓越。

表 1 不同方法性能对比结果

方法	性能指标		
	R@1	R@5	R@10
所提方法	0.985	0.974	0.969
文献 [1] 方法	0.847	0.788	0.802
文献 [2] 方法	0.789	0.871	0.814

4 结语

本文所提的基于深度文本聚类的关键词语义多模态挖掘方法,通过半监督学习算法精准抽取并融合多模态数据中的语义特征,有效过滤噪声并赋予与关键词语义挖掘任务相关的模态较高权值。得出结论如下:

(1) 在聚类效率方面,所提方法的 BWP 值稳步上升,短时间内即达到稳定状态,最高值达到了 0.9,显示出其快速且稳定的聚类性能。

(2) 在语义挖掘方面,所提方法的 R@1、R@5 和 R@10 指标分别高达 0.985、0.974 和 0.969,展现出卓越的关键词捕捉能力。

参考文献:

[1] 刘长红, 曾胜, 张斌, 等. 基于语义关系图的跨模态张量融合网络的图像文本检索 [J]. 计算机应用, 2022, 42(10): 3018-3024.

[2] ZHENG K F, WANG N, LIU J W, et al. An efficient multikey-word fuzzy ciphertext retrieval scheme based on distributed transmission for internet of things[EB/OL]. (2022-04-08) [2024-02-19]. <https://doi.org/10.1002/int.22886>.

[3] 杨帆, 宁博, 李怀清, 等. 基于语义增强特征融合的多模态图像检索模型 [J]. 浙江大学学报 (工学版), 2023, 57(2): 252-258.

[4] HUSSAIN S, ZIA M A, ARSHAD W. Additive deep feature optimization for semantic image retrieval[J]. Expert systems with applications, 2021, 170(5): 114545.

[5] 郜帅, 侯心迪, 刘宁春, 等. 多模态网络环境异构标识空间管控架构研究 [J]. 通信学报, 2022, 43(4): 26-35.

[6] 赵鹏, 马泰宇, 李毅, 等. 融合全模态自编码器和生成对抗机制的跨模态检索 [J]. 计算机辅助设计与图形学学报, 2021, 33(10): 1486-1494.

[7] 赵磊. 基于深度学习的多模态数据特征提取与选择方法研

究 [D]. 天津: 天津大学, 2015.

[8] 王晓莉. 基于差分进化算法的思政多模态语料库智能构建 [J]. 微型电脑应用, 2022, 38(5): 149-151.

[9] 尤博, 彭开香. 基于有效分类的多模态过程故障检测及应用 [J]. 上海应用技术学院学报 (自然科学版), 2015, 15(3): 242-247.

[10] 熊回香, 杨滋荣, 蒋武轩. 跨媒体知识图谱构建中多模态数据语义相关性研究 [J]. 情报理论与实践, 2019, 42(2): 13-18.

【作者简介】

周春良 (1996—), 男, 河南开封人, 硕士, 助教, 研究方向: 光学、图像识别。

杨畅畅 (1996—), 男, 河南商丘人, 硕士, 助教, 研究方向: 机器视觉、智能生产。

李晓辉 (1995—), 男, 河南平顶山人, 本科, 中级工程师, 研究方向: 计算机视觉、边缘计算。

(收稿日期: 2024-11-15)

(上接第 152 页)

一致, 当 $t > 300$ 时, 本文改进的算法逐渐体现出优势, 时间稍微增大的同时噪声检测率逐渐提高。而且, 经过改进的算法处理的新数据集得到的第一主成分包含的信息比传统算法承载更多的信息量, 因此可用较少的主成分提取更多的原始信息, 降维效果更显著, 改进的算法模型更具实用性。

4 结语

通过对机械零部件疲劳断裂数据采集及分析存在的问题进行梳理, 提出了一种基于加权模糊 C 均值 - 对数变换 PCA 算法用于疲劳断裂检测数据的分析优化处理。与传统的算法相比, 不仅能够提高噪声检测率, 而且第一主成分包含更多原始信息, 提高了数据降维效果, 不仅加强了试验数据的安全管理, 而且提高了疲劳断裂测试人员的管理能力。

参考文献:

[1] 王利东. 机械工程中金属材料的疲劳与断裂分析 [J]. 产品可靠性报告, 2024(6): 117-118.

[2] AMOOIE M A, LIJESH K P, MAHMOUDI A, et al. On the characteristics of fatigue fracture with rapid frequency

change[J]. Entropy, 2023, 25(6): 840.

[3] 郭尚志, 廖晓峰, 李刚, 等. 基于 PCA 的大数据降维应用 [J]. 计算机仿真, 2024, 41(5): 483-486.

[4] 孙家祥, 胡春玲. 方差参数和信噪比参数特定于父节点的全局耦合模型 [J]. 电脑知识与技术, 2023, 19(27): 9-12.

[5] CHEN T, KUO D, CHEN C Y Z. Retraction note: fuzzy c-means robust algorithm for nonlinear systems[J]. Soft computing, 2023, 3(28): 2769-2775.

[6] 汪杨海, 贺细平. 扩展欧几里德算法改进探讨 [J]. 电脑与信息技术, 2018, 26(6): 12-14.

【作者简介】

黄啸 (1986—), 通信作者 (email: weco_x@163.com), 男, 浙江宁波人, 硕士, 高级工程师, 研究方向: 材料力学性能表征。

叶序彬 (1980—), 男, 湖北大冶人, 硕士, 高级工程师, 研究方向: 材料力学性能表征。

(收稿日期: 2024-11-17)