

InSAR 时序形变数据的自监督对比学习聚类方法

吴晗飞¹, 俸彬¹, 李梦华^{1,2,3}, 杨梦诗⁴, 张震^{1,2,3}, 唐伯惠^{1,2,3}

1. 昆明理工大学 国土资源工程学院, 昆明 650093;

2. 云南省定量遥感重点实验室(筹), 昆明 650093;

3. 云南省山地灾害天空地一体化智慧监测国际联合实验室, 昆明 650093;

4. 云南大学 地球科学学院, 昆明 650500

摘要: 时序 InSAR TSInSAR (time-series Interferometric Synthetic Aperture Radar) 技术能够大范围获取形变, 成功应用于地质灾害监测、城市基础设施安全评价、矿区边坡监测等。然而, 时序 InSAR 技术获取的海量形变时间序列给形变场的精确解译带来了巨大的挑战。海量时间序列形变数据的自动分类, 对于精确解译形变信息、及时发现危险信号具有重要意义。本研究提出了一种基于自监督对比学习的 InSAR 时序形变深度聚类方法。该方法通过引入自监督对比学习框架, 增强模型在无标签数据上的聚类能力。同时, 针对时间序列数据增强方法在捕捉时间序列变换不变性方面的不足, 提出一种基于旋转形状一致性的数据增强策略。该策略通过对原始时间序列数据进行不同角度的旋转, 有效保持时间序列的形态相似性, 提升了聚类的准确性和鲁棒性。采用验证数据集, 将本研究提出的方法与 K -shape 方法在最佳聚类精度和归一化互信息值方面进行对比。验证结果发现: 本研究提出的方法分别较传统 K -means 方法提升了 25.8% 和 16.3%。说明本研究方法在捕捉时间序列特征与时间序列相似性度量方面表现出更好的准确性。此外, 本研究利用 2020 年 1 月—2022 年 10 月覆盖云南省个旧市卡房尾矿库的哨兵 1 号升轨数据集中所提取出的时间序列形变, 验证本研究所提出的方法。对该时间序列形变聚类分析之后得到了可靠的分类结果, 表明该方法可对 InSAR 形变时间序列进行有效分类。

关键词: 自监督对比学习, 数据增强, 时间序列形变, 形变聚类, 时间序列 InSAR

中图分类号: P23/P237

引用格式: 吴晗飞, 俸彬, 李梦华, 杨梦诗, 张震, 唐伯惠. 2025. InSAR 时序形变数据的自监督对比学习聚类方法. 遥感学报, 29 (7): 2442–2456

Wu H F, Feng B, Li M H, Yang M S, Zhang Z and Tang B H. 2025. Self-supervised contrastive learning clustering method for InSAR time series deformation data. National Remote Sensing Bulletin, 29 (7): 2442–2456 [DOI: 10.11834/jrs.20254393]

1 引言

随着时序 InSAR 技术的不断深入发展, 获取大范围长时序地表形变时间序列的能力显著提高。目前, 时序 InSAR 技术已成功应用在地质灾害监测 (Song 等, 2022; Liu 等, 2024)、城市基础设施安全评价 (Yang 等, 2024)、矿区边坡监测 (Hu 等, 2017) 等方面。然而, 在 InSAR 技术获取的海量时序形变数据中, 各类形变常常相互混杂, 难以对其进行解译并识别出危险信号。因此, 亟需发展对时序形变数据进行分类的相关技术以提高形

变的解译速度、降低解译难度, 帮助研究人员及时发现危险信号, 并做出反应。

时间序列被定义为按时间顺序排列的一组定量观测, 通常包括趋势成分、周期性成分以及随机成分共 3 个主要组成部分: (Kirchgässner 等, 2012)。对形变时间序列进行分类的最终目的是要从这些成分中识别出危险信号, 并作为指导防灾减灾的科学依据。Cigna 等 (2011) 使用简单线性回归模型对形变速率时间序列进行了分析, 随后又在此基础上, 使用预定义的变形模型计算出偏差指数 DI (Deviation Indexes) 来解释形变时间序

收稿日期: 2024-09-23; 预印本: 2025-01-19

基金项目: 国家自然科学基金(编号:42404036); 云南省基础研究计划(编号:202301AT070436)

第一作者简介: 吴晗飞, 研究方向为 InSAR 形变数据处理。E-mail: 20212201138@stu.kust.edu.cn

通信作者简介: 李梦华, 研究方向为 InSAR 地质灾害监测。E-mail: menghuali@kust.edu.cn

列 (Cigna 等, 2012); Berti 等 (2013) 提出一种基于条件序列统计检验的形变时间序列自动分类方法, 成功将地表形变分为了 6 类 (稳定、线性、二次、双线性、匀速不连续、变速不连续); Mirmazloumi 等 (2022a) 对该方法进行了改良, 进一步分类出了相位解缠误差 PUE (Phase Unwrapping Error)。

随着机器学习和深度学习的不断发展, 二者也逐渐被引入到形变时间序列的分类当中。首先是 Mirmazloumi 等 (2022b) 通过监督学习的方法, 有效地从形变数据中提取关键特征, 进而实现对地表形变类型的准确分类 (Kulshrestha 等, 2022)。在非监督分类方面, Van De Kerkhof 等 (2020) 通过使用 DBSCAN (Density-Based Spatial Clustering of Applications with Noise) 聚类方法识别空间密度相对较高的区域, 增强了 InSAR 结果的可解释性, 同样的 K -means 聚类方法被用于分类格林兰冰盖季节性变化 (Solgaard 等, 2022), Festa 等 (2023) 采用主成分分析 PCA (Principal Component Analysis) 来降低数据维度, Rygus 等 (2023) 利用均匀流形近似和投影 UNWAP (Uniform Manifold Approximation and Projection) 方法来达到相同的目的, 并应用聚类算法 HDBSCAN (Hierarchical Density-based Spatial Clustering of Applications with Noise) 和 K -means 对 InSAR 时间序列数据进行有关地表形变的聚类分析, Yang 等 (2024) 使用主成分分析 (PCA) 和 K -means 聚类方法对建筑环境下的形变时间序列进行了探索分析, 并讨论了 PCA 分解在 InSAR 形变解释中的意义。在深度学习方面, Li 等 (2024) 联合自堆叠编码器 SAE (Stacked AutoEncoder) 和卷积神经网络 CNN (convolutional neural network) 对昆明市地表沉降时间序列进行了分类, 并对分类结果做出了解释。

虽然现有研究已经取得了很多的成果, 但仍有诸多因素制约着 InSAR 形变时间序列的高效分类和解译。基于统计建模的传统方法以及有限的定义和假设, 使其对形变的表达能力受限, 难以揭示地表发生的复杂形变特征, 特别是难于捕捉对于一些具有危险性形变特征信号。监督分类时, 由于标签数据的获取成本高昂、样本不平衡、模型泛化应用能力的挑战等问题, 使得其应用与推广受到限制。非监督分类虽然无需制作样本数据,

但是捕捉 InSAR 时间序列的局部模式和时序依赖性方面存在一定的局限, 对定义和描述时间序列的相似性也面临挑战, 难以准确反映复杂的时间序列结构, 导致在复杂场景下的表现不够理想。

因此, 本研究拟结合 InSAR 时序数据的特点, 将自监督对比学习融入到深度聚类框架中, 以学习时间序列数据的内在结构相似性, 提出一种具有一定通用性的基于自监督对比学习的 InSAR 时间序列深度聚类方法; 同时通过在尾矿库区域的 InSAR 形变时间序列数据上开展实验和对比分析, 评估该方法在聚类精度、边界清晰度等方面的优势; 以期为 InSAR 时序数据的无监督特征提取与自动聚类提供更有效的技术手段, 助力地表形变监测、灾害识别与区域时空模式分析等相关应用研究。

2 研究方法

本研究将自监督对比学习的思想引入到时间序列的聚类当中, 从正负样本的角度定义时间序列相似性, 同时实现数据特征提取和聚类的联合优化。对比学习方法的核心思想是通过区分不同的实例数据, 利用数据增强生成的不同视图之间的相似性来进行学习。在特征空间中, 该方法将相似的样本聚集在一起, 同时将不相似的样本分离开, 以此优化分类的特征表示。

本研究构建了一个自监督对比学习 InSAR 时间序列深度聚类模型 (图 1)。该模型主要分为对比特征提取、对比约束与聚类分配共 3 个部分。在对比特征提取阶段, 首先利用数据增强技术, 在原始 InSAR 时间序列数据中生成正负样本对, 并比较两者的相似性, 学习样本的一般特征表示。在得到这些潜在特征空间后, 模型进入对比约束阶段, 进行实例级的对比学习, 以进一步细化特征表示。这不仅增强了特征空间的区分度, 也为聚类算法提供了一个更加分明的特征表示, 大大降低了聚类过程中的歧义性和不确定性。最后, 基于优化后的特征空间, 应用聚类算法, 将时间序列数据划分到不同的聚类簇中。通过这种自监督对比学习的方法, InSAR 时间序列聚类模型可以使用未标记数据分类地表形变时间序列。

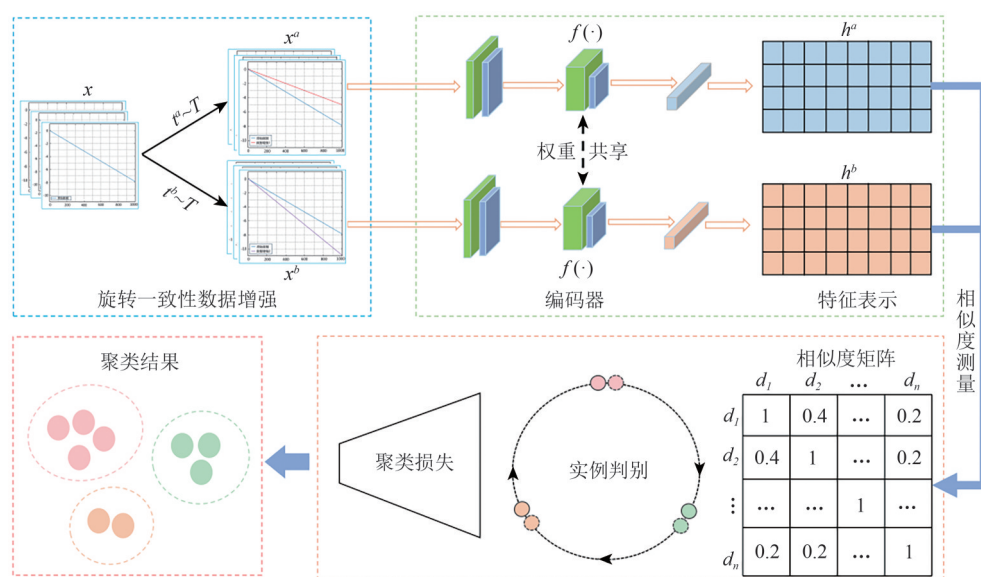


图1 基于对比学习的InSAR时间序列聚类模型示意图

Fig. 1 Schematic diagram of InSAR time series clustering model based on self-supervised contrastive learning

2.1 基于旋转形状一致性的数据增强方法

数据增强方法的好坏会直接影响下游任务的结果（杨博 等，2024）。为此，本研究结合对比学习在计算机视觉等领域的应用，并考虑 InSAR 时间序列数据的特点，提出了一种基于旋转形状一

致性的数据增强方法。该方法通过对原始数据进行不同角度的旋转，让有限的样本产生更多的价值，同时解决了现有时间序列数据增强方法无法充分捕捉时间序列变换不变性的问题。基于旋转形状一致性的数据增强方法如图2所示。

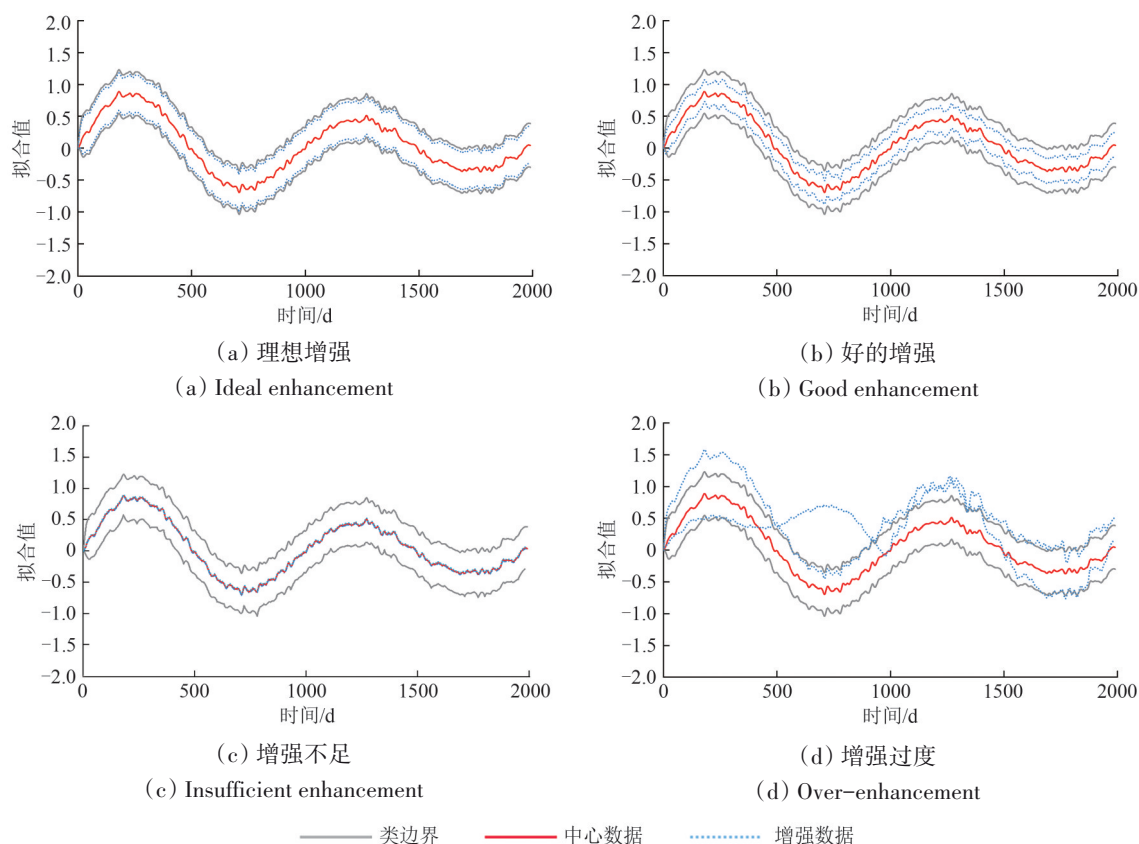


图2 基于旋转形状一致性的数据增强效果对比

Fig. 2 Data augmentation based on rotation shape consistency

旋转形状一致性是指经数据增强后的样本应尽可能保持在原始数据的区间内, 并且增强后的样本应尽量位于同一类别的区间内。本研究所提出的基于旋转形状一致性的数据增强方法具体计算公式如下:

$$x'(t) = x(t)\cos(\theta) - x(t-1)\sin(\theta) + g \quad (1)$$

式中, $x(t)$ 表示原始时间序列数据, $x'(t)$ 表示旋转后的时间序列数据, θ 为旋转角度, t 为时间, g 为高斯噪声。

2.2 时间序列特征提取

时序编码器 TSE (Time Series Encoder) 将时序数据转换为稀疏的中间层向量, 提取数据中的时序特征。选择合适的编码器对于准确提取数据的特征至关重要。本研究采用了 3 个一维卷积神经网络 (1D-CNN) 作为特征编码器, 该编码器结构由输入层、卷积层、激活层、批归一化 BN (Batch Normalization) 层以及 ReLU (Rectified Linear Unit) 激活函数等组成。输入层接收原始的时间序列数据; 接着, 数据通过卷积层进行特征提取, 每个卷积层后可能会跟随一个激活层和 BN 层来增加网络的非线性处理能力和稳定性; 最后, ReLU 激活函数被用于引入非线性特征并帮助网络学习复杂的数据表示。通过这一组合不仅捕获了 InSAR 时序数据的序列特征, 还通过稀疏表示降低了数据维度, 从而为后续的分析任务提供了高效且具有代表性的特征表示, 显著提高了数据处理的性能与效率。网络运算示意详见图 3。

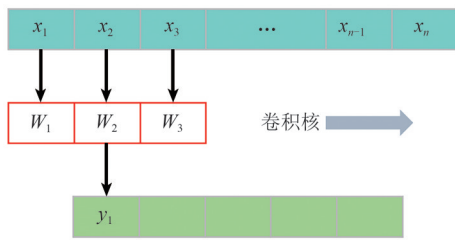


图3 1D-CNN 结构运算示意图

Fig. 3 Schematic diagram of 1D-CNN structure operation

对于原始的时间序列样本 $\mathbf{X} = \{x_i\}_{i=1}^n$, 其中所有的时序数据样本要求具有相同的步长 t 和相同的维数 d , 即 $x_i \in \mathbb{R}^{t \times d}$ 。原始时间序列数据 x_i 经过时序编码器后变成了高维时序特征向量 y_i , 如公式所示:

$$y_i = f_{\text{enc}}(x_i) \quad (2)$$

式中, $y_i \in \mathbb{R}^{t \times l}$, f_{enc} 为由 1D-CNN 构成的时序编码器, l 为编码器最后一层的隐藏层的状态数。

再经过由多层感知机 MLP (Multilayer Perceptron) 层组成的投影头 (Projection Head), 将高维特征向量表示矩阵投影至新的特征表示空间中。

$$\mathbf{O}_i^g = \text{ReLU}(\mathbf{W}_1^g \mathbf{R}^i + \mathbf{b}_1^g), t \in \{a, b\} \quad (3)$$

$$\mathbf{Z}^i = g(\mathbf{R}^i) = \mathbf{W}_2^g \mathbf{O}_i^g + \mathbf{b}_2^g \quad (4)$$

式中, $\mathbf{R}^i$ 是表示输入的特征向量, $t \in \{a, b\}$ 代表两个对比样本; $\mathbf{W}_1^g$, $\mathbf{W}_2^g$ 和 $\mathbf{b}_1^g$, $\mathbf{b}_2^g$ 代表 Projection Head 特征映射网络 $g(\cdot)$ 中的第一层和第二层的权重、偏置参数。 $\mathbf{O}_i^g$ 为第一层线性变换与 ReLU 激活后的中间表示。 $\mathbf{Z}^i$ 为最终投影得到的特征表示向量。

2.3 实例级对比约束

在对比特征提取阶段, 通过应用数据增强技术, 将 N 个原始数据扩展成 N 个二元组, 每个二元组内包含 2 个经过不同数据增强处理的样本, 同一数据不同的增强视为正样本对, 而来自不同数据的增强为负样本对。在一个 batch 中, 对于每个样本来说, 它将拥有一个正样本和 $2N-2$ 个负样本。根据所采用的数据增强方法, 这些样本分别被映射到 2 个不同的增强矩阵中, 投影结果 $\mathbf{Z}^i$ 分别得到 2 个特征表示矩阵 $\mathbf{Z}^a$ 和 $\mathbf{Z}^b$ 。

随后使用 $\text{sim}(\cdot)$ 表示余弦相似度, 对 $\mathbf{Z}$ 中的 2 个矩阵进行相似度测量, 如公式 (5) 所示。

$$\text{sim}(z_i, z_j) = \frac{\langle z_i, z_j \rangle}{\|z_i\| \cdot \|z_j\|} \quad (5)$$

式中, $\langle \cdot, \cdot \rangle$ 表示内积运算, $\|\cdot\|$ 为二范数。

对于给定任意一个特征表示 z_i, z_j , $(z_i, z_j) \in (\mathbf{Z}^a, \mathbf{Z}^b)$ 通过 NT-Xent (Normalized Temperature-scaled Cross Entropy) 损失函数计算余弦相似度, 最大化与正样本之间相似性, 并最小化 z_i 与其余 $2N-2$ 个负样本之间的相似性, 公式如下所示:

$$L_1 = -\log \frac{\exp(\text{sim}(z_i, z_j)/\tau)}{\sum_{k=1}^N l_{[k \neq i]} \exp(\text{sim}(z_i, z_j)/\tau)} \quad (6)$$

式中, z_i 和 z_j 分别代表同一个数据点的 2 个不同增强 (正样本对) 的特征表示。 $\text{sim}(z_i, z_j)$ 表示 z_i 和 z_j 之间的相似度。 N 是批次大小, 即每次计算损失时考虑的总样本数。 τ 是温度参数, 用于调节相似度分数的尺度。 $l_{[k \neq i]}$ 是一个指示函数, 当 $k \neq i$ 时取值为 1, 否则为 0。这确保了分母中不包括正样本

对 z_i 和自身的相似度。

2.4 软聚类分配损失

在实例级别对比约束获得样本特征表示之后，将软聚类分配损失（Soft Clustering Head Loss）引入到聚类过程中对其进行优化。软聚类分配损失是根据每一个数据点与各个聚类中心的距离赋予一个连续的概率分布，而非被硬性分配到某一单独聚类中。这种方法减轻了硬聚类错误可能带来的负面影响，并增强了聚类模型的鲁棒性。具体来说，该损失函数通过比较实际分配概率（即模型根据数据点与聚类中心的相似度计算得到的概率分布）与目标分配概率（基于最优化标准预设的理想概率分布）之间的差异来定义。

对于学习到的特征表达 $\mathbf{Z}$ ，我们可以通过应用 t 分布计算软标签如下：

$$q_{ij} = \frac{(1 + \|z_i - \mu_j\|^2 / \alpha)^{-\frac{\alpha+1}{2}}}{\sum_j (1 + \|z_i - \mu_j\|^2)^{-\frac{\alpha+1}{2}}} \quad (7)$$

式中， i 表示第 i 样本， u_j 表示第 j 个聚类中心， $z_i \in \mathbf{Z}$ ， q_{ij} 为样本 i 属于类 j 的概率，即软分配， α 通常取1。

软聚类分配损失通过学习高置信度分配来迭代地细化聚类分配，将软分配 Q 与辅助目标分布 P 相匹配， p_{ij} 计算公式如下：

$$p_{ij} = \frac{q_{ij}^2 / \sum_i q_{ij}}{\sum_j q_{ij}^2 / \sum_i q_{ij}} \quad (8)$$

将 q_{ij} 视为目标概率分布，而将模型预测的分配概率 p_{ij} 视为预测概率分布，则可以通过最小化2个分布之间的Kullback-Leibler (KL) 散度来训练模型：

$$L_2 = \text{KL}(P // Q) = \sum_i \sum_j p_{ij} \log \frac{p_{ij}}{q_{ij}} \quad (9)$$

这样模型的总体损失函数为

$$L = \lambda L_1 + (1 - \lambda) L_2 \quad (10)$$

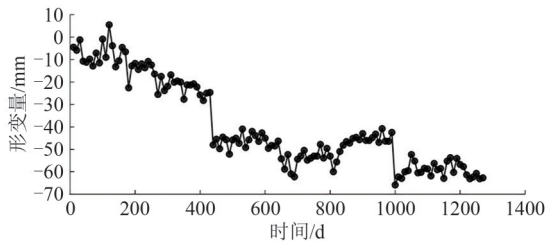
式中， λ 为权重系数。

3 基于真实标记样本数据集的聚类模型对比分析

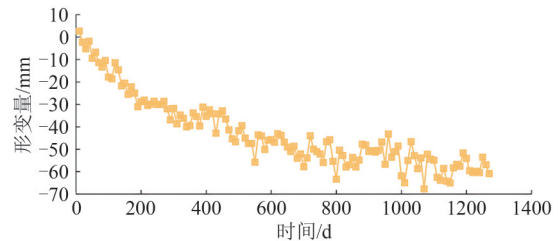
为展示模型性能，本研究进行了多项实验活动，包括对比实验，用于对比不同算法的性能；消融实验，用于评估算法中各部分的贡献；以及参数敏感性实验，探讨不同参数设置对结果的影响。文章选取了 K -means ((Krishna和Murty, 1999)、 K -shape (Paparizos 和 Gravano, 2015)、Agglomerative clustering (Murtagh 和 Contreras, 2012)、DAE (Denoising Autoencoder Algorithm) (Vincent 等, 2010)、DEC (Deep Embedded Clustering) (Xie 等, 2016)、DTC (Deep Temporal Clustering) (Madiraju 等, 2018) 等6种方法进行对比，其涵盖了3种传统聚类算法和3种基于深度学习的聚类方法，可较为全面的对模型进行对比评估。

3.1 数据集及评价指标

本研究将时序形变信号分为稳定型、线性型、加速型、减速型以及相位解缠误差型PUE (Phase Unwrapping Error) 等5个类别。在真实时序InSAR结果中进行标注，共筛选标注了5000个数据样本，分别对应5种不同的地表形变类型，每一类别分别包含了1000个样本实例 (Li 等, 2024)，用于模型的验证。图4展示了各类时序样本在时间-形变量维度下的代表性曲线特征，以突出不同类别之间的变化趋势差异，具体分布如图4所示。



(a) PUE 型
(a) PUE type



(b) 减速型
(b) Deceleration type

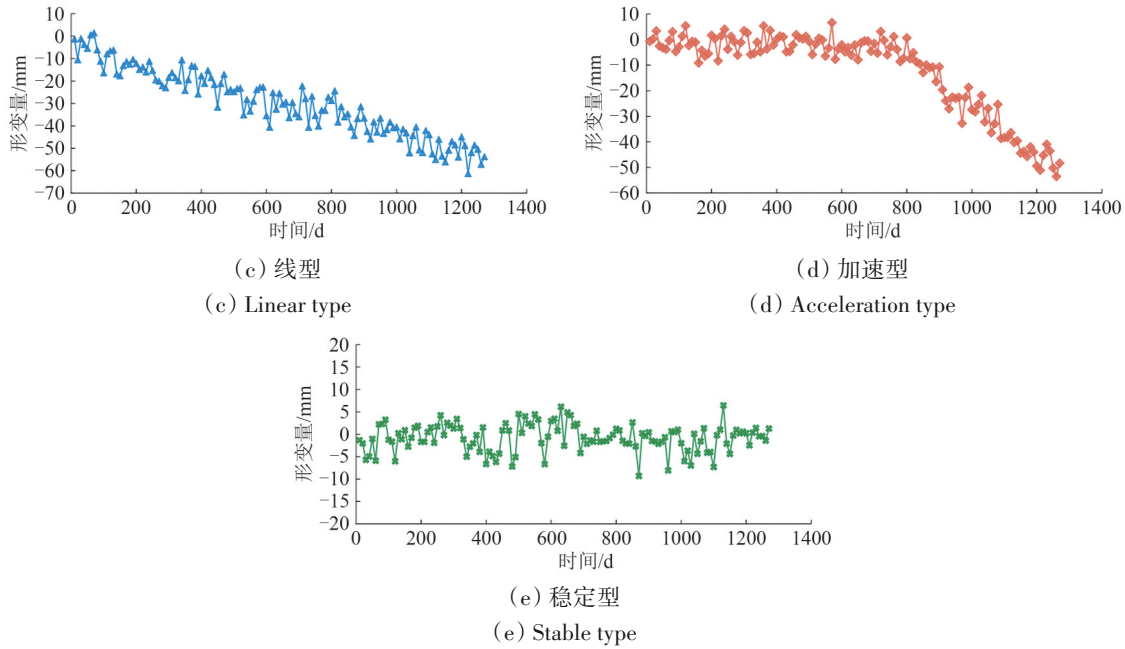


图4 5种不同形变趋势类型定义

Fig. 4 Five types of deformation trend type definition style diagram

为了衡量不同聚类算法的效果,本研究使用了聚类精度 ACC (Clustering Accuracy) (Yang 等, 2010) 和归一化互信息 NMI (Normalized Mutual Information) (Xu 等, 2003) 进行聚类性能检验。ACC 是通过比较真实标签与预测标签之间的匹配度来计算的。NMI 是衡量 2 个标签集相似度的标准化度量。计算公式分别如下:

$$ACC = \max_m \frac{\sum_{i=1}^N l\{l_i = m(c_i)\}}{N} \quad (11)$$

式中, l_i 表示第 i 个数据的真实值, c_i 表示预测值, $m(\cdot)$ 表示数据点的预测值和真实值之间的最优映射, N 为样本总数, $l\{\}$ 是指示函数, 条件成立是返回 1, 否则返回 0。

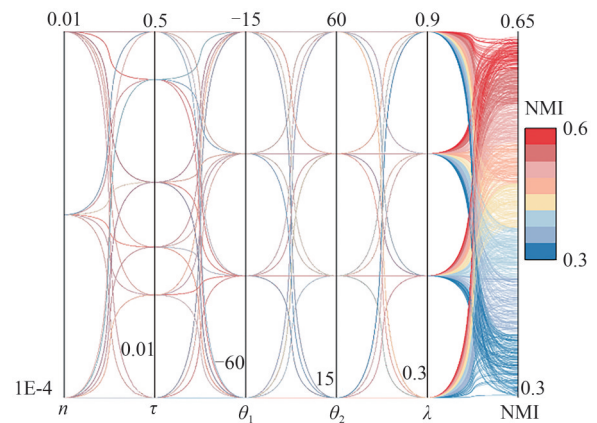
$$NMI = \frac{2I(l_i; c_i)}{H(l_i) + H(c_i)} \quad (12)$$

式中, $I = (l_i; c_i)$ 表示所有数据点的真实标签 $l = \{l_1, l_2, \dots, l_N\}$ 和预测聚类分配 $C = \{cl_1, cl_2, \dots, cl_N\}$ 之间的互信息。 $H(\cdot)$ 表示熵函数。ACC 和 NMI 范围均在区间 $[0, 1]$ 内, 值越接近于 1, 表示聚类的性能越好。

3.2 模型参数分析及确定

为使所提出的模型取得最优聚类效果, 本研究使用组合调参方法对模型中的主要参数进行了评估与改进, 包括学习率 η , 温度参数 τ , 数据增

强中旋转角度参数 θ_1 、 θ_2 , 以及模型的权重参数 λ 。在前期的单独参数实验中发现, 将每个参数约束在一个对应的区间范围上可以取得较好的试验结果, 因此设置 $\eta \in \{0.01, 0.001, 0.0001\}$, $\tau \in \{0.01, 0.03, 0.3, 0.5, 0.7\}$, $\theta_1 \in \{-60^\circ, -45^\circ, -30^\circ, -15^\circ\}$, $\theta_2 \in \{15^\circ, 35^\circ, 45^\circ, 60^\circ\}$, $\lambda \in \{0.3, 0.5, 0.7, 0.9\}$, 使用组合调参方法可以得到 1152 种参数组合形式, 对调参后的实验结果进行可视化分析, 可视化结果如图 5 所示。

图5 网络调参可视化结果(其中 η 、 τ 、 θ_1 、 θ_2 、 λ 分别表示学习率、温度参数、旋转角度参数 1 和 2、模型的权重参数)Fig. 5 Network parameter adjustment visualization results (Among them η , τ , θ_1 , θ_2 , λ respectively represent the learning rate, temperature parameter, rotation angle parameters 1 and 2, and the model weight parameter)

对 NMI 结果进行回溯读取, 可以发现在聚类性能表现较好的参数组中的各参数取值较为相近, 验证了调参过程中各项参数取值的稳定性, 并最终确定了最佳参数组合为: $\eta = 0.001$, $\theta_1 = -15^\circ$, $\theta_2 = 30^\circ$, $\tau = 0.05$, $\lambda = 0.7$, 初始 batch 大小为 256。后续的聚类分析实验均采用该参数配置, 并在构建的数据集上进行共 400 轮的训练。整个算法流程如下所述:

对输入的 InSAR 时间序列数据进行处理。每轮训练迭代包含以下 6 个主要步骤:

(1) 划分 batch: 将训练集划分为小批量子集, batch 大小为 256。

(2) 数据增强: 对每个样本应用旋转变换等方式生成两个不同视图, 形成正样本对, 见式 (1)。

(3) 特征提取: 通过编码器对增强样本进行表征学习, 输出高维特征表示矩阵, 见式 (2)。

(4) 非线性映射: 将提取的特征向量输入至投影头 (Projection Head) 中, 获得最终特征表示 $\mathbf{Z} = [\mathbf{Z}^a, \mathbf{Z}^b]$ 。

(5) 损失函数计算: 依次计算以下 3 种损失:

实例级对比损失 (式 (6));

聚类分配损失 (式 (9));

总体损失函数 (式 (10))。

(6) 参数优化: 使用 Adam 优化器对网络参数进行更新, 完成一次训练迭代。

3.3 对比聚类模型性能分析

3.3.1 对比试验分析

将提出的方法在数据集上与其他多个传统算法进行了比较, 对比结果见图 6。本研究中提出的深度时间对比聚类方法 DTCC (Deep Temporal Contrastive Clustering) 在所有方法中表现最佳, ACC 达到 0.78, NMI 为 0.64, 高于其他聚类方法。这说明 DTCC 通过对比学习的架构, 在实例级别的对比约束与相似性判别的基础上捕捉了时间序列数据的关键特征, 并且在聚类任务上取得了更准确和一致的结果。其中: K-means 作为一种应用最广泛的聚类算法之一, 在时间序列数据上的表现相对较弱, ACC 和 NMI 分别为 0.39 和 0.33, 这可能是由于其简化的聚类假设, 不适应时间序列数据的复杂性; AC (凝聚聚类) 的表现较差, ACC 为 0.14, NMI 为 0.31, 可能是由于其在构建聚类时依赖于距离度量, 而时间序列之间的相似度不易准确度量; 针对时间序列设计提出的 K-shape 算法, 通过捕获数据形态特征表现出突出的效果, 达到了 0.62 的 ACC 和 0.44 的 NMI, 在 3 种传统的聚类算法中表现最佳, 原因在于其是依据形状进行时间序列相似性判别, 对几种典型地表形变数据集的区分更为容易。

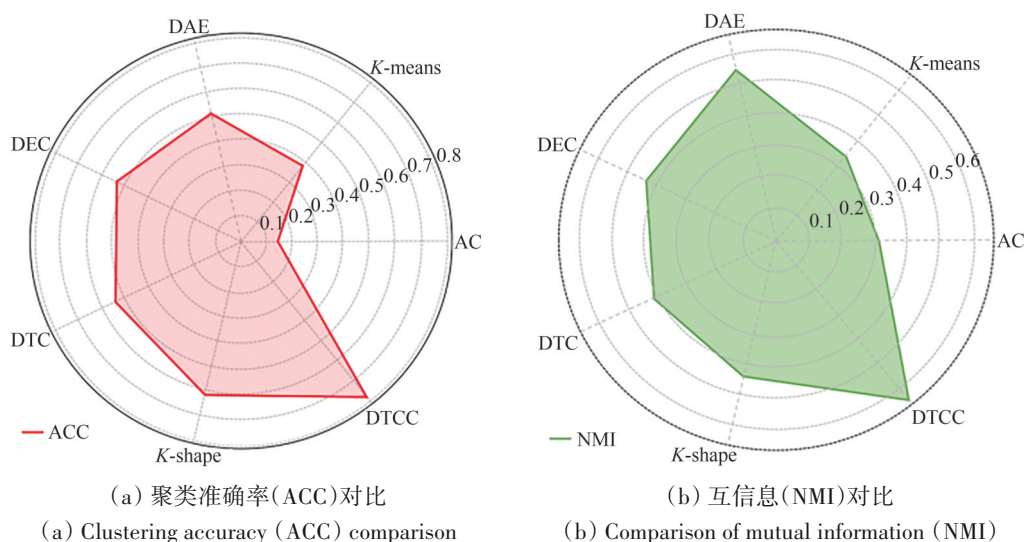


图 6 聚类方法性能对比图

Fig. 6 Radar chart comparison of clustering method performance

DAE 以及 DEC 通过深度学习技术来提取特征, 展现了较为平衡的性能。其中: DAE 的 NMI 高达

0.55, 表明其在提取内在特征上的有效性; DEC 联合了特征学习与聚类优化, 其 ACC 和 NMI 分别为

0.53 和 0.44。表明 DAE 和 DEC 在聚类任务上各有侧重与提升,但还是可以有进一步的优化空间。在本研究中所使用的深度时间聚类 DTC (Deep Temporal Clustering) 方法,将深度学习应用于时间序列聚类中,得到了 0.54 的 ACC 和 0.42 的 NMI,验证了所使用的时间序列编码器在 InSAR 时间序列任务中的时序特征提取能力,能够满足时序聚类下游任务的需要 (图 6)。

3.3.2 消融实验分析

为验证所提算法的有效性,本研究基于 1D CNN 的深度聚类算法进行消融实验比较。在本次消融试验中,本研究逐步引入不同的模型组件,并观察到了模型性能的逐渐提高,详见表 1。可见:作为基础模型的 DTC,其单独应用时的 ACC 和 NMI 分别为 0.54 和 0.42;当加入了对比学习模块 (Contrastive Learning) 之后,ACC 提升到了 0.62, NMI 也增加到了 0.52,说明对比学习思想的引入对模型性能有积极正面影响;随后,在 DTC+CC 的基础上加入了多层感知器 (MLP),得到的 DTC+MLP+CC 模型进一步提高了 ACC 到 0.74 和 NMI 到 0.57,表明 MLP 投影头的使用有助于避免编码器生成的特征空间因对比损失而产生扭曲,从而增强其在提取时间序列数据特征方面的效果;最终,通过进一步加入软聚类分配损失模块 SCH (Soft Clustering Head Loss),ACC 得到了进一步提升,达到了 0.78, NMI 达到了 0.64,这一增长验证了 SCH 模块在分配聚类标签上性能提升的有效性。整体而言,这些实验结果突出了每一步新增的网络结构对于提升模型在具体任务上性能的重要性。

表 1 消融实验结果

Table 1 Ablation experiment results		
方法	ACC	NMI
DTC	0.54	0.42
DTC+CC	0.62	0.52
DTC+MLP+CC	0.74	0.57
DTC+MLP+CC+SCH	0.78	0.64

3.3.3 数据增强方法分析

本研究使用基于旋转形状一致性的时间序列数据增强方法对数据进行增强,为了检验该方法

的有效性,进行了一系列试验,对比了添加噪声、序列翻转、缩放和序列切片等不同数据增强技术在同一数据集上的表现,试验结果见表 2。

表 2 数据增强方法对比实验结果

Table 2 Comparative experimental results of data enhancement methods

方法	ACC	NMI
高斯噪声	0.66	0.50
序列翻转	0.59	0.42
缩放	0.67	0.53
序列切片	0.55	0.31
旋转	0.78	0.64

可见与传统数据增强方法相比,本研究提出的基于旋转形状一致性的数据增强方法在对比学习中展现出了显著的性能优势。虽然添加高斯噪声和序列缩放对性能有所提升,但这种优势并不十分明显,尤其是序列切片在 NMI 上的表现甚至更差,导致性能下降。因此,本研究介绍的形状驱动的数据增强技术通过利用序列的信号强度来进行形状的扩展,这一方法不仅维护了数据的初始特征,同时也提高了每个样本的信息内容,从而增强了模型处理时间序列数据的性能与精度。

4 基于实测数据的深度对比聚类模型验证分析

4.1 试验区概况

云南省红河州个旧市的矿山行业历史悠久。该地区矿产资源丰富,是全国最大的锡金属生产地,也是世界上最早开始生产锡的基地 (Wu 等, 2022)。卡房尾矿库是一座平地型的三等尾矿库,地理位置位于滇中褶皱带、扬子台地与思茅地块的交汇区域。这一地区地质特征复杂,存在多条断层并且在这些断层的交接处岩石结构破碎,受到地质活动的影响较大。该区域的降水量较高,沿断层带常见溶蚀漏斗和垂直溶洞的发展,表明水文地质条件活跃。如图 7 所示。尾矿库矿区面积为 91.3 万 m²,设计库容为 1853 万 m³,截至 2020 年,该尾矿库累积存放了约 1000 万 t 尾矿,这些尾矿根据其性质可以分为尾粉砂、尾粉土、尾粉质粘土和尾黏土层 (彭沛宇和王卫华, 2020)。

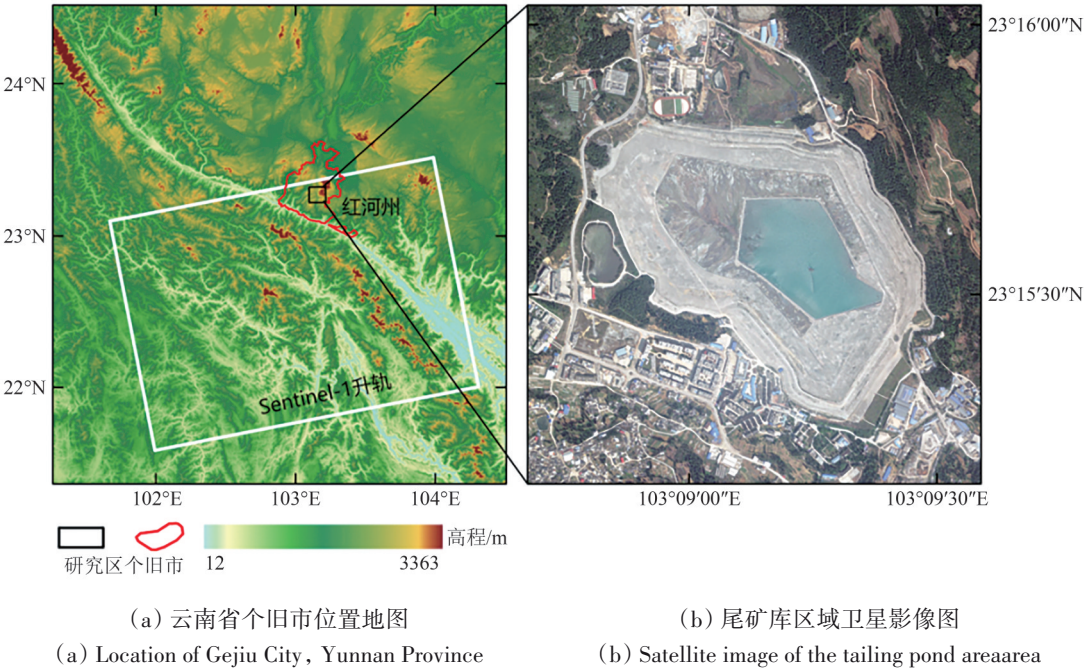


图7 研究区尾矿库的地理位置分布
Fig. 7 Geographical distribution of tailings ponds in the study

4.2 试验数据及平台

本研究选择2020年1月—2022年10月的时段、覆盖试验区的84幅Sentinel-1升轨SAR影像进行验证。在数据处理阶段，采用山区精度表现较好的AW3D30 DEM数据来消除地形相位的影响(Li等, 2023)。此外，为了提高影像配准和基线估计的精度，进行了必要的轨道校正。

本研究所有算法代码都是基于64位Python 3.7编程语言和Tensorflow2.4深度学习框架开发的。本试验使用的计算机配置为Intel (R) Core (TM) i9-12900H 2.70 GHz 显卡，6 GB 显存。GPU 运算平台为NVIDIA 公司CUDA11.0，深度学习GPU加速库为CUDNN8.0。

4.3 试验区时序InSAR形变结果

本研究利用SBAS-InSAR技术对个旧市卡房尾矿库坝区展开验证，利用覆盖该地区的升轨Sentinel-1雷达影像数据提取了该地区地表形变信号。通过详细的分析，成功获取了卡房尾矿库坝区共6966个相干点的地表年平均形变速率。图8展示了尾矿库坝及其周边区域的地表位移速率情况，其中坝体中心区域沿南北向展现出显著的红色集中区，该集中区向外围逐渐扩散，最大形变速率达到-19.7 mm/a。

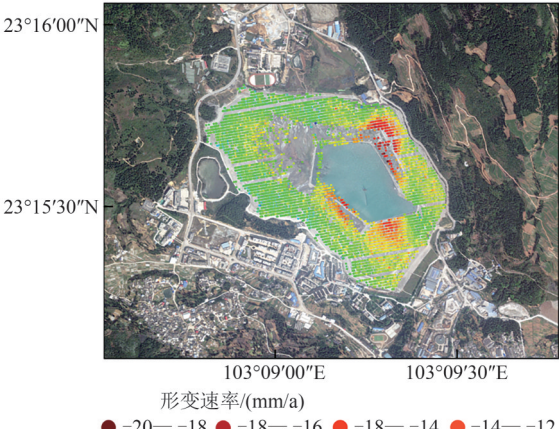


图8 卡房尾矿库坝视线向形变速率监测结果
(以整景影像的平均速率为参考点)
Fig. 8 Monitoring results of line-of-sight deformation rate of Kafang tailings pond (With the average rate of the whole scene image as reference point)

4.4 深度对比聚类模型识别分析

本研究为分析不同聚类方法在卡房尾矿坝形变时序识别中的表现，图9展示了利用DTCC、K-means与K-shape方法对卡房尾矿坝区域6966个PS点进行聚类分析后在不同聚类数(K=2与K=3)下的聚类均值时间序列图。结果显示，使用3种方

法得到的时间序列均值变化趋势存在明显差异, 征方面具有不同的能力。这表明不同聚类算法在揭示尾矿坝区域内形变特

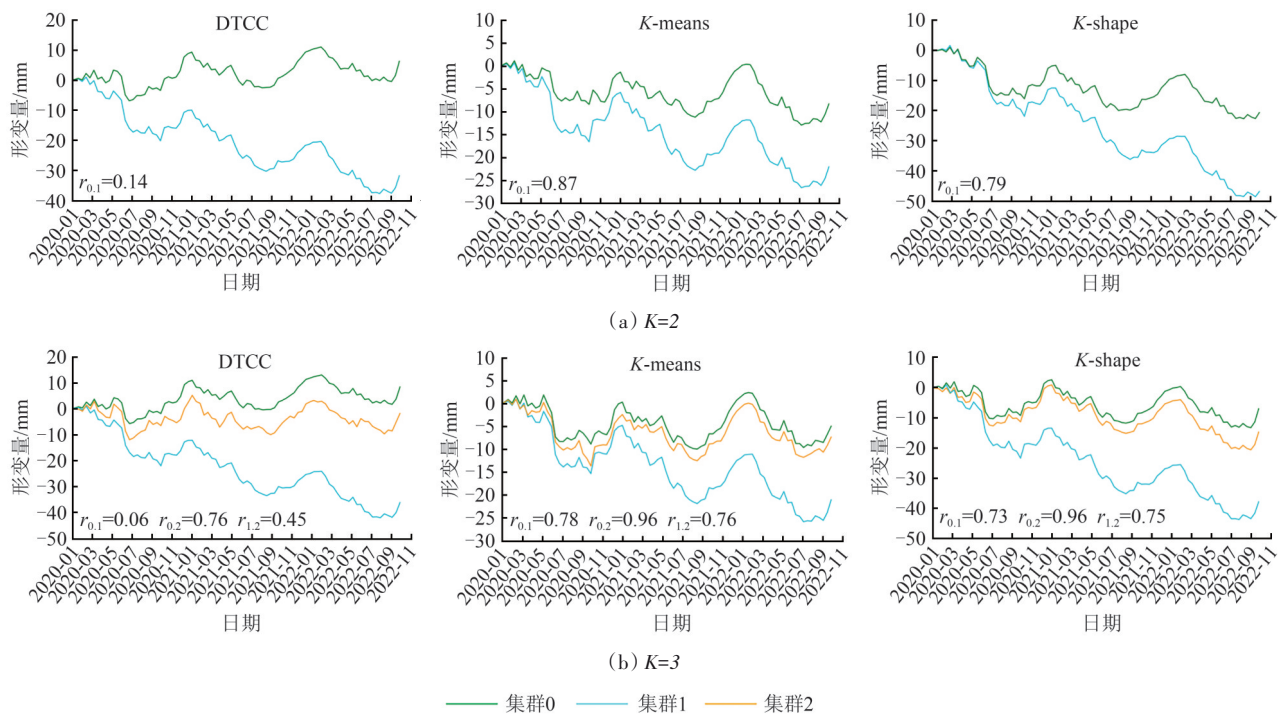


图9 卡房尾矿库在DTCC与K-means、K-shape方法下聚类类别均值的时间序列图(曲线展示了各算法在不同聚类数下的类别均值时间序列趋势,标注的 r_{ij} ,表示类别间的相关系数,数值越小表示聚类边界越清晰)

Fig. 9 The time series diagram of the clustering category mean of the Kafang tailings pond under DTCC, K-means, and K-shape methods (The curve shows the time series trend of the category mean of each algorithm under different cluster numbers. The correlation coefficient r_{ij} between categories is also marked. The smaller the value, the clearer the cluster boundary)

由图9(a)可见,在聚类簇数 $K=2$ 的条件下,DTCC方法能够有效划分出两类趋势显著不同的形变模式:其中集群0表现为相对平稳的周期性波动,集群1呈现出明显的线性下降叠加振幅变化的趋势,二者之间的相关系数仅为0.14,显示出良好的类间区分度,表明DTCC方法能够清晰地区分不同类型的形变特征,具备较好的类间区分能力。相较之下,K-means得到的两类形变序列差异不大,均为周期性波动,相关系数高达0.87,说明其难以有效识别差异较大的形变类别;K-shape方法虽然在一定程度上区分了周期幅度差异,相关性为0.79,但聚类边界仍不如DTCC清晰。

由图9(b)可见,在聚类簇数 $K=3$ 设置下,DTCC进一步细化了类别结构:集群0延续周期性波动特征,集群1呈下降趋势,新增的集群2虽也呈周期性变化,但与集群0的相关性为0.76,仍具有一定差异性。整体上,DTCC在 $K=3$ 时依然能较好地保持类间差异性,进一步验证了其在区分不

同形变模式上的有效性。而K-means方法聚类结果间相似性仍高,识别出的3个类别之间的形变特征依然较为相似,集群0与集群2的相关性高达0.96,整体相关性均在0.78以上,显示其在边界识别方面仍存在局限。K-shape方法在细节上能够刻画振幅差异,但各类别之间的相似性较高。集群0和集群1仍然表现出平稳的周期性波动,而新增的集群2则展现出较大的振幅变化。集群0与集群2的相关系数达到0.96,类别重叠问题依然存在。总体程度上,DTCC方法在形变特征分类上展现了较好的区分能力。当簇的个数划分为3时,DTCC能够保持类间的差异性,有效地区分不同的形变模式。而K-means方法在类间区分上表现不佳,尤其是在识别形变边界方面存在局限。K-shape方法在低聚类数时具备较好的区分能力,但在较高聚类数下其区分效果不如DTCC。

综上,DTCC方法在两种聚类数设置下均展现出较好的形变类别分离能力和时序结构表达效果,

优于 K -means 与 K -shape, 特别是在形变模式混合复杂、趋势差异显著的场景中优势突出。

图 10 展示了三种聚类方法 (K -means、 K -shape 和 DTCC) 在不同聚类数 ($K=2$ 与 $K=3$) 条件下的方差统计结果, 用于评估各方法对形变时序数据的分类集中度与类别边界清晰度。图 10 (a) 中, 聚类簇数被设置为 2 时, 3 种聚类方法 (K -

means、 K -shape 和 DTCC) 在集群点数和方差上表现出显著差异。DTCC 在集群 0 中的方差最低 (0.11), 聚类效果最集中, K -shape 次之 (0.18), K -means 方差最高 (0.22)。在集群 1 中, DTCC 和 K -shape 方法 (方差值分别为 0.45 和 0.51) 的表现优于 K -means (方差值为 0.58), 显示其对不同模式形变的识别更为清晰稳定。

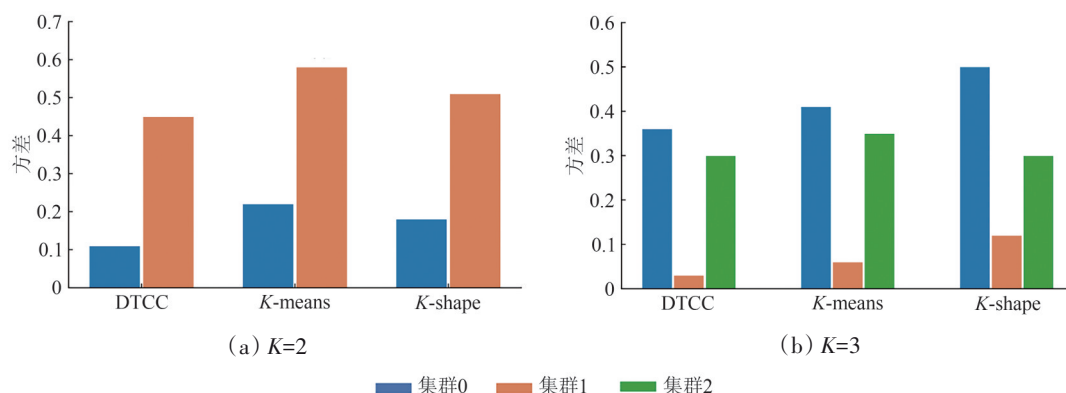


图 10 DTCC 与 K -means、 K -shape 方法下的聚类方差统计图 (出于显示的目的, 方差经过标准化统计)
Fig. 10 DTCC, K -means, and K -shape cluster variance statistics (For display purposes, variances are standardized)

图 10 (b) 中, 聚类簇数被设置为 3 时, 3 种聚类方法 (K -means、 K -shape 和 DTCC) 在集群特征上表现出不同的效果: DTCC 在 3 个集群中均表现出较低方差, 尤其在集群 1 中方差显著小于其他方法 (最低值为 0.11), 说明其在多簇条件下依然保持良好的聚类性能。相比之下, K -means 方法在集群 1 中的方差较高, 而 K -shape 则在多个类簇之间分布较均衡, 但整体方差高于 DTCC。DTCC 在低方差聚类方面的效果最佳, 突出了 DTCC 方法在区分时间序列数据相似性和确定聚类类型边界上的优势。综合比较可见, DTCC 在多个聚类数条件下均表现出更低的类内方差和更合理的样本分布, 表明其聚类结果更集中、边界更清晰。

图 11 展示了使用两种不同的聚类算法对数据进行 t -SNE (t 分布随机邻域嵌入) 降维可视化的聚类结果。每一行代表一种聚类方法, 每一列代表不同的聚类数 K (2 和 3)。颜色梯度表示不同集群的分布情况, 从中可以观察到各方法在数据空间中的聚类效果差异。集群的个数为 2 时, 图 11 (a) DTCC 中结果显示为两个明显的区域, 蓝色和红色的区域较为清晰而且紧凑, 界限分明。这表明使用 DTCC 方法的聚类能较好地识别数据中的非线性

结构和复杂特征。 K -means 与 K -shape 虽然也区分出了两个主要的聚类, 但存在边界模糊和部分重叠的情况。

随着 K 增加到 3, DTCC 进一步细化了聚类, 图 11 (b) DTCC 展示了蓝色和灰色两个主要簇, 以及一个类别明显的红色簇。蓝色簇显示为较集中和一致的群体, 而红色簇则表现为清晰划分的小集群, 这种聚类布局有助于快速识别和分析数据中的主要组成部分。而 K -means 中的虽然区分了红色、灰色和蓝色簇, 但这种聚类的结果簇间不够紧凑, 只是粗略的进行了数据的表达, 可能忽略了细微的数据特征和结构。 K -shape 则表现出集群内部分布不均与混乱的问题, 原因可能在于 K -shape 注重时间序列的相位对齐, 对于复杂的形变数据, K -shape 在分割集群时容易出现模糊的边界和集群内分布不均的情况, 在降维后的平面空间中, 难以形成清晰的边界。

整体来看, DTCC 方法展现了最清晰的边界, 能够有效区分不同集群, 其在处理分布不均和具有明显群体边界的数据集时的特定优势, 便于快速把握 InSAR 时序数据的关键特征, 突出重要的形变信息。

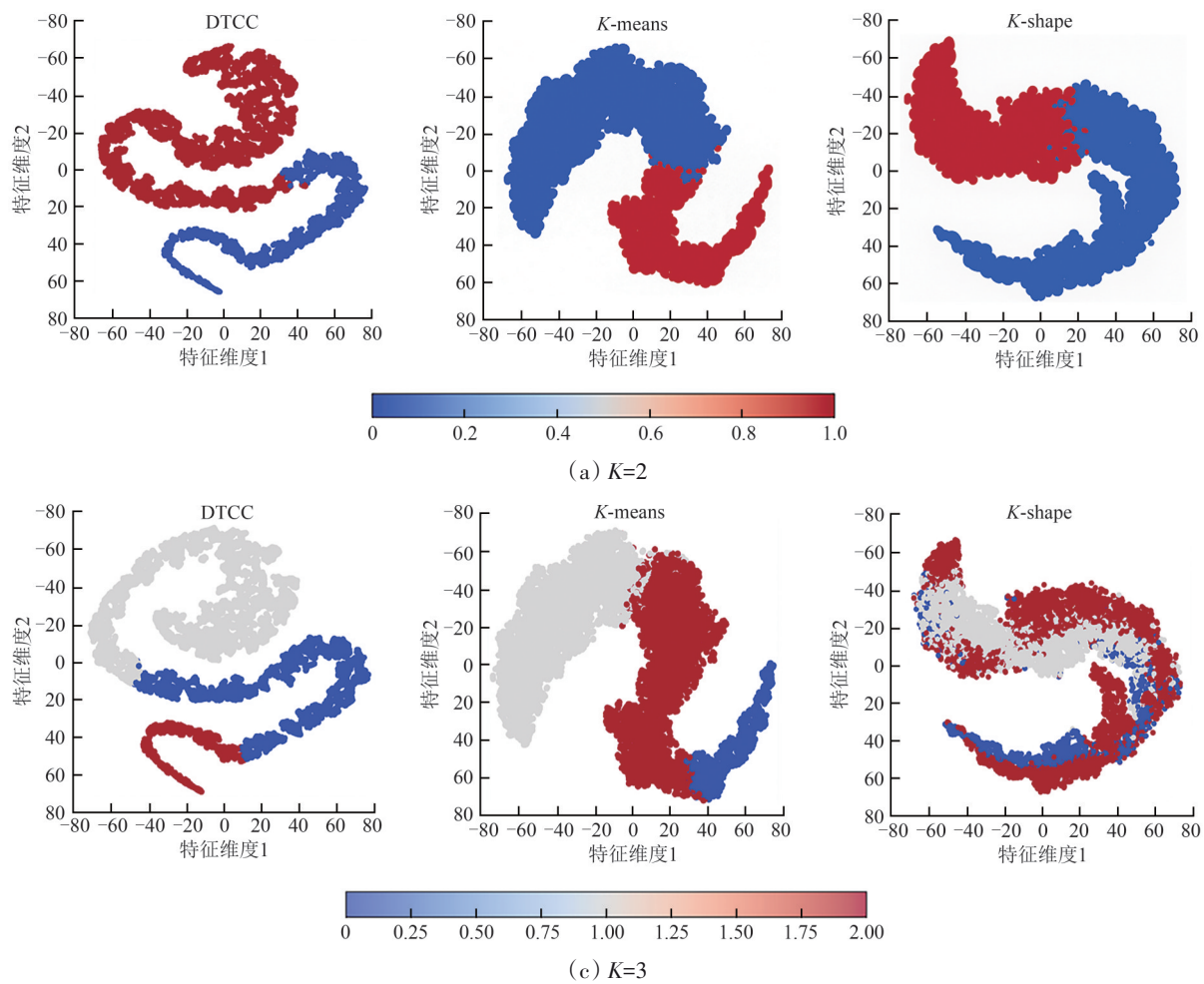
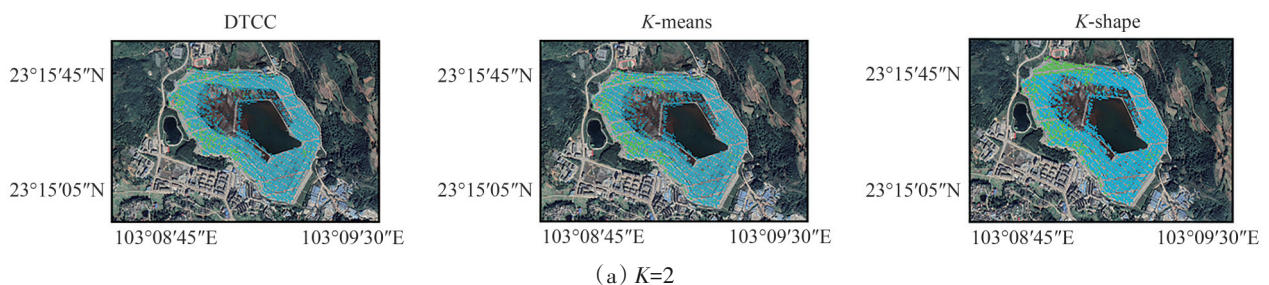


图 11 DTCC, K -means, and K -shape 3 种方法的降维可视化结果 (色标条表示聚类编号仅用于区分类别, 无实际单位)
Fig. 11 Dimensionality reduction visualization results of DTCC, K -means, and K -shape methods (Colorbar indicates cluster ID, used for visual separation only, with no physical unit)

图 12 中展示了不同聚类方法 (DTCC、 K -means 和 K -shape) 在坝体形变数据上的空间分布效果。DTCC 方法在不同簇个数的情况下均表现出清晰的集群分布, 在 $K=3$ 时能够细致地划分出不同形变区域, 准确反映坝体形变的空

间分布效果。相比之下, K -means 和 K -shape 方法在聚类边界的清晰度和形变区域的空间分布上略显不足。我们将所构建的模型用于卡房尾矿库的 InSAR 形变时间序列分类, 结果显示本研究提出的 DTCC 方法聚类结果类别间边界清晰, 分类结果相对聚集。这证明了本研究提出的模型对 InSAR 形变时间序列分类的有效性。



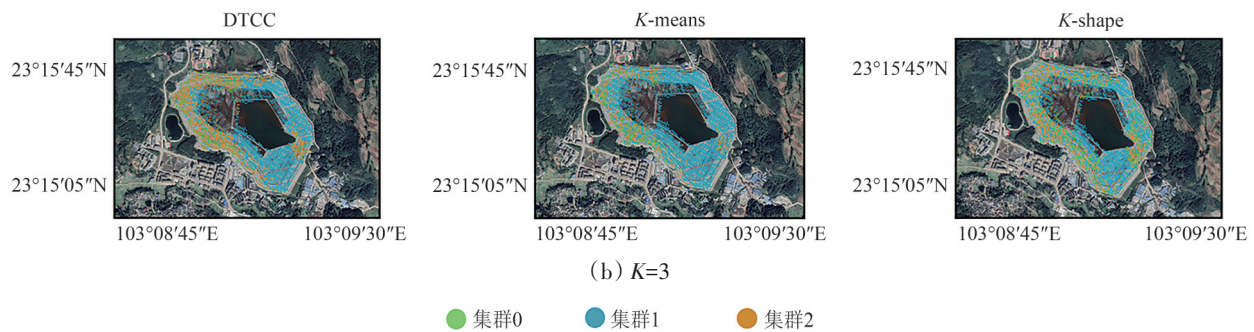


图12 尾矿库坝体时序形变数据在3种不同方法下的空间聚类效果对比

Fig. 12 Comparison of spatial clustering effects of tailings dam deformation data under three different methods

5 结 论

为提升 InSAR 时序数据在无监督条件下的聚类精度与时序结构识别能力,本研究以典型地表形变序列为研究对象,采用自监督对比学习方法,提出了一种结合对比损失与聚类分配优化的深度聚类模型 (DTCC),并与传统聚类方法如 K -means 与 K -shape 进行对比分析与性能验证。主要结论如下:

(1) 通过基于旋转形状一致性的数据增强方法扩充数据集,构建正负样本对,从而在无需显式标签的环境中定义时间序相似性。与传统的数据增强方法相比,该方法有效引入了形状结构约束,上增强了对形变模式内在结构的识别能力。

(2) DTCC 在真实地表形变数据集上的聚类结果优于对比方法,其中聚类一致性指标 NMI 最高达到 0.78,聚类精度 ACC 达到 0.64,相较 K -means (NMI 0.33, ACC 0.38) 与 K -shape (NMI 0.44, ACC 0.61) 有明显提升。

(3) 在卡方尾矿库坝体的时序形变案例识别中,验证了所提方法在真实复杂场景中的适应能力与判别性能。在聚类数 $K=2$ 设置下,DTCC 所得两类均值时间序列的相关系数仅为 0.14,低于 K -means (0.87) 与 K -shape (0.79),表明其在类别可分性和边界判别方面表现更优,验证了模型对复杂时序结构的识别分类能力,具备在大尺度 InSAR 监测与地质灾害识别中的实际应用潜力。

参考文献 (References)

Berti M, Corsini A, Franceschini S and Iannacone J P. 2013. Automated classification of persistent scatterers interferometry time series. *Natural Hazards and Earth System Sciences*, 13(8): 1945-1958

[DOI: 10.5194/nhess-13-1945-2013]

Cigna F, Del Ventisette C, Liguori V and Casagli N. 2011. Advanced radar-interpretation of InSAR time series for mapping and characterization of geological processes. *Natural Hazards and Earth System Sciences*, 11(3): 865-881 [DOI: 10.5194/nhess-11-865-2011]

Cigna F, Tapete D and Casagli N. 2012. Semi-automated extraction of Deviation Indexes (DI) from satellite Persistent Scatterers time series: tests on sedimentary volcanism and tectonically-induced motions. *Nonlinear Processes in Geophysics*, 19(6): 643-655 [DOI: 10.5194/npg-19-643-2012]

Festa D, Novellino A, Hussain E, Bateson L, Casagli N, Confuorto P, Del Soldato M and Raspini F. 2023. Unsupervised detection of InSAR time series patterns based on PCA and K-means clustering. *International Journal of Applied Earth Observation and Geoinformation*, 118: 103276 [DOI: 10.1016/j.jag.2023.103276]

Hu X, Oommen T, Lu Z, Wang T and Kim J W. 2017. Consolidation settlement of Salt Lake County tailings impoundment revealed by time-series InSAR observations from multiple radar satellites. *Remote Sensing of Environment*, 202: 199-209 [DOI: 10.1016/j.rse.2017.05.023]

Kirchgässner G, Wolters J, Hassler U. Introduction to modern time series analysis [M]. Springer Science & Business Media, 2012.

Krishna K and Murty M N. 1999. Genetic K-means algorithm. *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, 29(3): 433-439 [DOI: 10.1109/3477.764879]

Kulshrestha A, Chang L and Stein A. 2022. Use of LSTM for sinkhole related anomaly detection and classification of InSAR deformation time series. *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, 15: 4559-4570 [DOI: 10.1109/JSTARS.2022.3180994]

Li M H, Wu H F, Yang M S, Huang C and Tang B H. 2024. Trend classification of InSAR displacement time series using SAE-CNN. *Remote Sensing*, 16(1): 54 [DOI: 10.3390/rs16010054]

Li M H, Yin X B, Tang B H and Yang M S. 2023. Accuracy assessment of high-resolution globally available open-source DEMs using ICESat/GLAS over mountainous areas, a case study in Yunnan Province, China. *Remote Sensing*, 15(7): 1952 [DOI: 10.

- 3390/rs15071952]
- Liu X J, Zhao C Y, Yin Y P, Tomás R, Zhang J, Zhang Q, Wei Y J, Wang M and Lopez-Sanchez J M. 2024. Refined InSAR method for mapping and classification of active landslides in a high mountain region: Deqin County, southern Tibet Plateau, China. *Remote Sensing of Environment*, 304: 114030 [DOI: 10.1016/j.rse.2024.114030]
- Sai Madiraju N, Sadat S M, Fisher D, et al. Deep temporal clustering: Fully unsupervised learning of time-domain features[J]. *arXiv e-prints*, 2018: arXiv: 1802.01059.
- Mirmazloumi S M, Gambin A F, Palamà R, Crosetto M, Wassie Y, Navarro J A, Barra A and Monserrat O. 2022a. Supervised machine learning algorithms for ground motion time series classification from InSAR data. *Remote Sensing*, 14(15): 3821 [DOI: 10.3390/rs14153821]
- Mirmazloumi S M, Wassie Y, Navarro J A, Palamà R, Krishnakumar V, Barra A, Cuevas-González M, Crosetto M and Monserrat O. 2022b. Classification of ground deformation using sentinel-1 persistent scatterer interferometry time series. *GIScience and Remote Sensing*, 59(1): 374-392 [DOI: 10.1080/15481603.2022.2030535]
- Murtagh F and Contreras P. 2012. Algorithms for hierarchical clustering: an overview. *WIREs: Data Mining and Knowledge Discovery*, 2(1): 86-97 [DOI: 10.1002/widm.53]
- Paparrizos J and Gravano L. 2015. k-shape: efficient and accurate clustering of time series//*Proceedings of the 2015 ACM SIGMOD International Conference on Management of Data*. Melbourne, Victoria: ACM: 1855-1870 [DOI: 10.1145/2723372.2737793]
- Peng P Y and Wang W H. 2020. Effects of soil heavy metal compounds on the micro-aggregate stability in Kafang tailing pond. *Agricultural Research in the Arid Areas*, 38(4): 215-220, 244 (彭沛宇, 王卫华. 2020. 卡房尾矿库土壤重金属化合物对微团聚体稳定性的影响. *干旱地区农业研究*, 38(4): 215-220, 244) [DOI: 10.7606/j.issn.1000-7601.2020.04.27]
- Rygu M, Novellino A, Hussain E, Syafudin F, Andreas H and Meisina C. 2023. A clustering approach for the analysis of InSAR time Series: application to the Bandung Basin (Indonesia). *Remote Sensing*, 15(15): 3776 [DOI: 10.3390/rs15153776]
- Solgaard A M, Rapp D, Noël B P Y and Hvidberg C S. 2022. Seasonal patterns of greenland ice velocity from sentinel-1 SAR data linked to runoff. *Geophysical Research Letters*, 49(24): e2022GL100343 [DOI: 10.1029/2022GL100343]
- Song C, Yu C, Li Z H, Utili S, Frattini P, Crosta G and Peng J B. 2022. Triggering and recovery of earthquake accelerated landslides in Central Italy revealed by satellite radar observations. *Nature Communications*, 13(1): 7278 [DOI: 10.1038/s41467-022-35035-5]
- Van De Kerkhof B, Pankratius V, Chang L, Van Swol R and Hanssen R F. 2020. Individual scatterer model learning for satellite interferometry. *IEEE Transactions on Geoscience and Remote Sensing*, 58(2): 1273-1280 [DOI: 10.1109/TGRS.2019.2945370]
- Vincent P, Larochelle H, Lajoie I, Bengio Y and Manzagol P A. 2010. Stacked denoising autoencoders: learning useful representations in a deep network with a local denoising criterion. *The Journal of Machine Learning Research*, 11: 3371-3408 [DOI: 10.5555/1756006.1953039]
- Wu H, Zheng X Y, Fan H D and Tian Z M. 2022. Deformation monitoring of tailings reservoir based on polarimetric time series InSAR: example of Kafang Tailings Reservoir, China. *Remote Sensing*, 14(15): 3655 [DOI: 10.3390/rs14153655]
- Xie J Y, Girshick R and Farhadi A. 2016. Unsupervised deep embedding for clustering analysis//*Proceedings of the 33rd International Conference on Machine Learning*. New York: JMLR.org: 478-487 [DOI: 10.5555/3045390.3045442]
- Xu W, Liu X and Gong Y H. 2003. Document clustering based on non-negative matrix factorization//*Proceedings of the 26th Annual International ACM SIGIR Conference on Research and Development in Informaion Retrieval*. Toronto: ACM: 267-273 [DOI: 10.1145/860435.860485]
- Yang B, Luo J C, Song Y T, Wu H T and Peng F R. 2024. Time series clustering method based on contrastive learning. *Computer Science*, 51(2): 63-72 (杨博, 罗嘉琛, 宋艳涛, 吴宏涛, 彭甫镕. 2024. 基于对比学习的时间序列聚类方法. *计算机科学*, 51(2): 63-72) [DOI: 10.11896/jsjx.221200038]
- Yang M S, Li M H, Huang C, Zhang R S and Liu R. 2024. Exploring the InSAR deformation series using unsupervised learning in a built environment. *Remote Sensing*, 16(8): 1375 [DOI: 10.3390/rs16081375]
- Yang Y, Xu D, Nie F P, Yan S C and Zhuang Y T. 2010. Image clustering using local discriminant models and global integration. *IEEE Transactions on Image Processing*, 19(10): 2761-2773 [DOI: 10.1109/TIP.2010.2049235]

Self-supervised contrastive learning clustering method for InSAR time series deformation data

WU Hanfei¹, FENG Bin¹, LI Menghua^{1,2,3}, YANG Mengshi⁴, ZHANG Zhen^{1,2,3}, TANG Bohui^{1,2,3}

1. School of Land and Resource Engineering, Kunming University of Science and Technology, Kunming 650093, China;

2. Yunnan Provincial Key Laboratory of Quantitative Remote Sensing (Under Preparation), Kunming 650093, China;

3. Yunnan International Joint Laboratory for Integrated Intelligent Monitoring of Mountain Hazards by Sky and Land, Kunming 650093, China;

4. School of Earth Sciences, Yunnan University, Kunming 650500, China

Abstract: The time-series Interferometric Synthetic Aperture Radar (InSAR) technique is widely recognized for its capability to monitor large-scale deformations, making it instrumental in applications such as geologic disaster monitoring, urban infrastructure safety assessment, and mining slope evaluation. This study aims to address the challenges in deciphering large-scale deformation time-series data obtained through InSAR. By leveraging self-supervised contrastive learning, we enhance the classification and clustering capabilities of InSAR time-series deformation data, which is critical for identifying geohazard signals and supporting infrastructure safety evaluation.

We propose a novel deep clustering framework based on self-supervised contrastive learning to enhance clustering performance on unlabeled deformation time-series data. To overcome the limitations of traditional time-series data augmentation techniques, a rotational consistency-based data augmentation strategy is introduced. This strategy maintains morphological similarity by rotating the original time series at different angles, enabling the model to better capture invariance in time-series transformations. The method's clustering performance is evaluated against traditional K-means clustering, with key metrics such as clustering accuracy and normalized mutual information (NMI) used for comparison. The framework is further validated using deformation data extracted from the Sentinel-1 ascending orbit dataset, covering the Kafang tailings pond in Gejiu City, Yunnan Province from January 2020 to October 2022.

The proposed method outperformed traditional K-means clustering, achieving a 25.8% improvement in clustering accuracy and a 16.3% increase in NMI. Compared with the K-shape method, the proposed method shows better accuracy in capturing time-series features and similarity measurement. The clustering analysis performed on the Sentinel-1 dataset successfully classified deformation time series into meaningful groups, revealing distinct deformation patterns and effectively identifying potential danger signals.

The integration of self-supervised contrastive learning with InSAR time-series analysis enhances the interpretability and classification of deformation patterns. The proposed method provides a robust and efficient tool for geohazard monitoring, urban infrastructure evaluation, and mining slope safety assessment. This approach has potential for broader applications in large-scale remote sensing data analysis.

Key words: self-supervised contrastive learning, data enhancement, time series analysis, deformation clustering, time-series InSAR

Supported by National Natural Science Foundation of China (No. 42404036); Yunnan Provincial Basic Research Program (No. 202301AT070436)