

Application of Voformer-EC Clustering Algorithm to Stock Multivariate Time Series Data

1st Ning Xin

School of AI and Advanced Computing
Xi'an Jiaotong-Liverpool University
Suzhou, China
Ning.Xin21@student.xjtlu.edu.cn

2nd Shaheen Khatoon

School of Arch., Comp. and Engineering
University of East London
London, United Kingdom
s.khatoon@uel.ac.uk

3rd Md Maruf Hasan

School of AI and Advanced Computing
Xi'an Jiaotong-Liverpool University
Suzhou, China
MdMaruf.Hasan@xjtlu.edu.cn

Abstract—Clustering stocks with similar increasing and decreasing trends has always been a challenging problem. Despite the extensive research on stock forecasting, the balance between effective clustering and speed remains challenging. Traditional multivariate time series clustering methods are difficult to guarantee high speed with high accuracy. This study proposes the Voformer-EC model as a new approach to solve this issue, enhancing the analysis of stock-related multivariate time series data. The Voformer-EC model takes both time features and volatility and utilises the Voformer neural network to extract time features and implement clustering. The data recorded every 60 minutes of the Nifty 50 Index from February 2nd to February 28th in 2015 was applied to the traditional model and compared with the Voformer-EC model. The results showed that the Voformer-EC model was significantly better than the traditional model. Follow-up studies consider applying the Voformer-EC model to temperature and precipitation to identify drought-prone areas and implement specific risk mitigation strategies in a targeted manner.

Index Terms—Voformer-EC Neural Network, Multivariate Time Series Clustering, Volatility Activation Function

I. INTRODUCTION

Significant advancements in deep learning and machine learning have catalyzed growing research interest in time series analysis among scholars and experts. Additionally, the field of time series clustering has undergone extensive developments [1]- [3], and its applications have been employed across various domains, including life expectancy prediction in insurance [4], carbon neutrality tracking [5], media data analysis, seasonal analysis [2], and financial data modelling [6]. Time series data exhibit inherent properties such as trends, seasonal effects, cyclical fluctuations, and residuals. However, traditional statistical models like Error Trend Seasonality [7] and Auto-Regressive Integrated Moving Average (ARIMA) [8] have well-documented limitations in time series analysis [9]- [10]. This renders the development of efficient and accurate clustering algorithms for high-dimensional, multivariate time series problems an ongoing challenge.

Time series clustering approaches can be categorized based on the clustering object into whole time series, subsequence, and time point clustering [3]. Typically, clustering methods comprise two key stages: distance measures and cluster algorithm [1]. For spatial time series, clustering techniques can be grouped into hierarchical, partitioning-based, and overlapping

methods according to the treatment of the spatial dimension [11]. Traditional clustering algorithms include partition-based, density-based, and hierarchical categories [12]. Various techniques can be derived by combining different clustering algorithms with distance measures. However, the choice of distance metric tends to be more impactful than the clustering method itself, as distance measures will greatly affect the clustering results. Commonly used distance measures include Euclidean Distance (ED) [13], shape-based distances [14], and Dynamic Time Warping (DTW) [15]. ED is the most widespread among these, while DTW has garnered substantial research focus in recent years [15]- [16]. DTW employs time warping to determine the optimal alignment between two time series, enabling shape-based similarity assessment. However, DTW has high time complexity and, similar to ED, is sensitive to outliers and noise as it considers all time points [16].

Neural networks for time series generate a feature similarity matrix as a distance measure is a core idea of this research. The core innovation lies in the distance measure and clustering method compared to traditional approaches. Specifically, this study employs the variant Informer model to extract time series features. The model's specificity for time series data is further enhanced by implementing the Volatility Activation Function (VAF). By considering both time and shape features of time series, the proposed approach substantially improves computational efficiency, reducing time complexity to $O(L \cdot \log_L)$ [17] compared to traditional distance measures. After feature extraction, a clustering algorithm is applied to group the time series data based on the learned feature distances. By combining specialized neural networks with clustering, this approach achieves state-of-the-art performance for time series clustering.

This study utilizes the proposed Voformer-EC model to analyze the Nifty 50 stock index over the period of February 2nd to 28th, 2015. The model generates a feature similarity distance matrix and performs clustering on the multivariate time series data. Comparative benchmarking is conducted against traditional time series clustering algorithms.

The critical contributions of this work are as follows:

1. A novel multivariate time series clustering neural network, Voformer-EC, is proposed. It performs better than

traditional clustering methods while greatly reducing time consumption.

2. The model can rapidly cluster stocks with similar price trends. This enables more informed decision-making by investors and financial managers.
3. The work advances interdisciplinary research at the intersection of machine learning and finance. The proposed techniques readily apply to fields such as agriculture and climate analytics.

II. METHODOLOGY

A. Basic models

1) *Informer*: The Informer model is a deep learning architecture for time series forecasting proposed recently [18]. Unlike Convolutional Neural Networks (CNNs) and Recurrent Neural Networks (RNNs), Informer utilizes a ProbSparse self-attention mechanism instead of standard self-attention. This is designed to capture temporal patterns and dependencies in the input time series data.

Specifically, ProbSparse self-attention employs a query sparsity measurement $M(q_i, K)$ to identify and focus on the most relevant key-query interactions. This measurement comprises two terms: 1) the Log-Sum-Exp (LSE) of query q_i over all keys, representing the aggregated attention weights, and 2) the arithmetic mean over all keys, acting as a normalization factor. A higher sparsity score $M(q_i, K)$ indicates the query has a more "informative" attention distribution over the keys. This allows ProbSparse self-attention to concentrate modeling capacity on the dominant long-range dependencies in lengthy time series.

ProbSparse self-attention improves Informer's computational efficiency compared to the standard Transformer through selective attention [19]. Specifically, it restricts each key to attend to only the top- μ dominant queries, where μ is proportional to the logarithm of the query sequence length L_Q . This is implemented via:

$$A(Q, K, V) = \text{Softmax}\left(\frac{\bar{Q}K^T}{\sqrt{d}}\right)V \quad (1)$$

Where $\bar{Q}$ is a sparse matrix containing just the Top- μ queries from Q . The softmax weighted attention is thus focused on the most relevant key-query interactions. By adjusting the sparsity hyperparameter μ , ProbSparse self-attention achieves much greater efficiency than Transformer and LSTM models for multivariate time series. Moreover, Informer provides flexibility to customize various hyperparameters, optimizing performance for different applications.

2) *Volatility Activation Function (VAF)*: The Volatility Activation Function (VAF) captures input data volatility to improve model accuracy, especially for time series analysis [20]. The VAF is defined as:

$$\text{Volatility} = \sqrt{\frac{\sum (P_{mean} - P_i)^2}{n}} \quad (2)$$

Where P_{mean} is the mean input, P_i are the individual input values, and n is the number of inputs. With $P_i = x_i$ and $P_{mean} = \bar{x}$, this reduces to:

$$f(x_1, \dots, x_n) = \sqrt{\frac{\sum (\bar{x} - x_i)^2}{n}} \quad (3)$$

By simplifying the notation and deriving the equation, we get:

$$f(x_1, \dots, x_n) = \sqrt{\frac{\sum (\bar{x} - x_i)^2}{n}} = \frac{1}{\sqrt{n}} \cdot \sqrt{\sum (\bar{x} - x_i)^2} \quad (4)$$

The volatility slope indicates changes in variance, allowing VAF to capture time series dynamics adaptively. This improves modelling of temporal structure compared to static activation functions.

3) *Density Clustering Model*: The Extreme Clustering (EC) [21] algorithm is a variant of Density Peaks Clustering (DPC) [22] that also utilizes density-based clustering. EC enhances DPC with several key improvements to facilitate the identification and separation of clusters and noise points:

- An extreme-searching procedure to locate cluster centers, rather than just density peaks. This captures more representative cluster cores.
- A saddle point detection method to identify cluster edges. This enables more defined cluster boundaries.
- An efficient nearest-neighbor graph construction algorithm. This improves cluster connectivity modeling.
- Explicit categorization of points as cores, edges, and noise. This benefits cluster interpretation.

EC improves the clustering quality compared to standard Density Peaks Clustering. EC demonstrates strong empirical performance on complex data distributions.

B. Voformer-EC

The proposed Voformer-EC model integrates Informer networks, Volatility Activation Function, and Extreme Clustering to create a specialized multivariate time series clustering framework. Compared to traditional algorithms, Voformer-EC strengthens long-range sequence modelling and feature extraction for time series data. This significantly improves efficiency and accuracy. Fig 6a is a schematic diagram of the model.

Additionally, Voformer-EC inherits strengths from both its components. The Voformer backbone enables robust time series feature capture and noise resilience. Meanwhile, Extreme Clustering provides flexibility to identify arbitrary cluster shapes.

Moreover, Voformer-EC offers adaptability to different applications through modular components and tunable parameters. Overall, this integrated model provides state-of-the-art performance on the challenging task of multivariate time series clustering.

Voformer-EC focuses on effectively capturing time series volatility and shape features for clustering. The main process is that the encoder takes multivariate time series data as input and generates a high-dimensional feature representation.

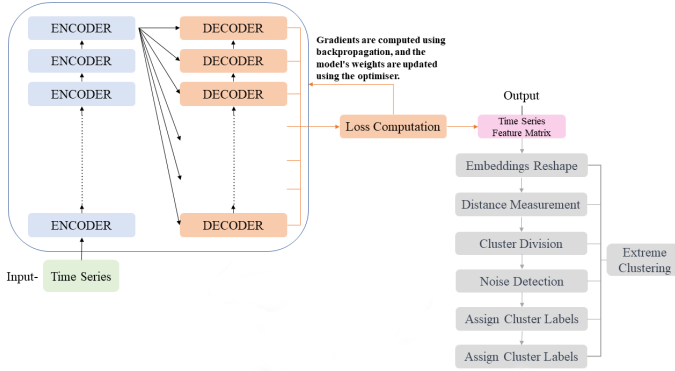


Fig. 1: Voformer-EC Model Schematic Diagram

The Decoder role is to reconstruct the original input from the encoded features. This reconstruction output is a high-dimensional feature similarity matrix used for density-based clustering. This encoding emphasizes time series feature extraction using a modified Informer architecture. The encoded representation is then clustered with the Extreme Clustering algorithm.

III. DATA DISPLAY AND EXPERIMENTS

A. Data Display

This study utilizes time series data from the Nifty 50 index recorded at 60-minute intervals from February 2nd to 28th, 2015. The aim is to cluster stocks exhibiting similar price trends over this period. The overarching goal is developing an effective solution for large-scale multivariate time series clustering that can generalize across disciplines.

The raw dataset comprises the closing price of 50 Indian stocks across February 2015, obtained from the Kaggle repository [23]. Given the scale of this data, the analysis focuses on the contiguous period from February 2nd to 28th for 52 selected indexes. Extraneous features were removed during preprocessing to derive a focused dataset for modelling. The closing price time series for these 52 stocks is visualized in Fig 6b, where each line represents a distinct stock index.

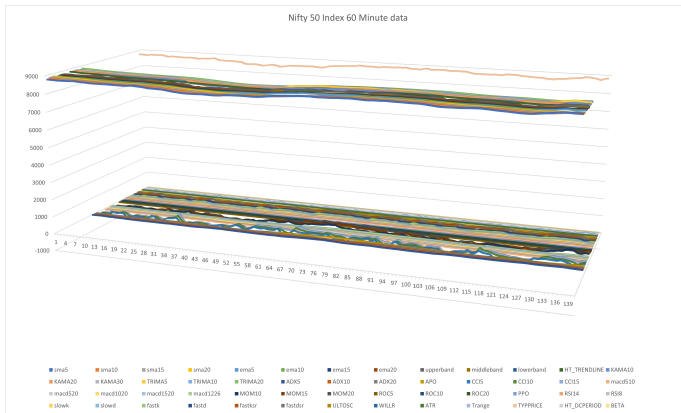


Fig. 2: Nifty 50 Index 60 Minute data

Fig 6b visualizes the multivariate time series dataset, where the x-axis represents time, the y-axis indexes each stock, and the z-axis shows the current price. Qualitatively, the data exhibit some key properties. First, it presents a bipolarity, with stock prices constrained within two extreme value ranges overall. Second, all indexes demonstrate volatility over the month, with some stocks showing strong fluctuations. Quantitatively analyzing this multivariate dataset's temporal patterns, correlations, and volatility will enable effective clustering of similarly-behaving assets.

B. Clustering Methods

Clustering techniques comprise two key components, the clustering algorithm and distance measure, significantly influencing performance. This study benchmarks three widely adopted distance metrics: Euclidean, Dynamic Time Warping (DTW), and K-Shape. We focus comparative analysis on established traditional methods for clustering algorithms, including partition-based, density-based, and hierarchical categories [1]. While numerous novel algorithms exist, these foundational approaches are chosen due to their proven robustness across diverse datasets and time series tasks [24]. The proposed Voformer-EC model will compare with the combinations of traditional time series clustering techniques by evaluation indexes.

TABLE I: Composition overview of contrastive clustering methods

Clustering Method	Distance Measure	Category
K-means	Euclidean, DTW, Shape-based	Partitional
DBSCAN	Euclidean, DTW	Density-based
Agglomerative	Euclidean, DTW	Hierarchical

To evaluate the performance of Voformer-EC, comparative benchmarking is undertaken against traditional clustering algorithms paired with established distance measures. As summarized in Table I, the combinations analyzed include k-means, DBSCAN, agglomerative clustering, and their variants incorporating Euclidean distance, Dynamic Time Warping, and K-Shape distances. After completing this comparative benchmarking, internal clustering validation metrics are systematically applied to quantify the quality of the resulting clusters. This analysis aims to situate Voformer-EC within the landscape of classical clustering techniques for time series data.

C. Evaluation Methods

In the absence of ground truth labels, internal clustering validation metrics provide practical quantitative estimates of result quality. This study employs five common metrics:

- **Within-cluster Sum of Squares (inertia):** Quantifies within-cluster dispersion around centroids. Lower values indicate tighter clusters.
- **Silhouette Score:** Evaluates cohesion of points within their cluster versus the nearest neighbouring cluster.

Ranges -1 to 1, with higher values corresponding to superior clustering.

- **Davies-Bouldin Index:** Measures the ratio of within-cluster scatter to between-cluster separation. Lower values suggest improved clustering.
- **Calinski-Harabasz Index:** Assesses the ratio of between-cluster to within-cluster dispersion. Higher values imply better-defined clusters.
- **Dunn Index:** Computes the ratio of minimum inter-cluster to maximum intra-cluster distance. Higher values correspond to dense, well-separated clusters. However, sensitive to outliers.

These metrics aim to quantify cluster cohesion, separation, and validity to evaluate relative clustering performance.

D. Clustering Experiments

The cluster analysis process for test data is divided into two main stages: the various combinations load data to cluster, followed by comparative data clustering effects under various methods.

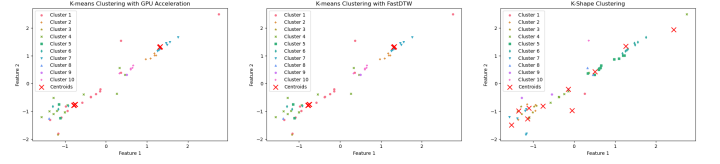
As a foundational clustering technique, k-means is benchmarked in combination with different distance measures. Despite its popularity, k-means has limitations in handling multivariate data due to sensitivity to outliers and difficulty modelling complex cluster shapes. This study combines k-means with three established distance metrics with Euclidean, Dynamic Time Warping (FastDTW), and K-Shape to cluster the Nifty 50 Index dataset. The efficacy of each k-means variant is quantified through the five internal validation metrics. This analysis provides baseline clustering performance to compare with other models.

Considering the substantial data volume, GPU acceleration is employed for the k-means algorithm. Subsequently, FastDTW replaces the default Euclidean distance metric, chosen for its linear time complexity. Finally, based on the shape-based concept, the K-shape is employed, and Clustering is performed with k ranging from 2 to 20. The results are visualized in Fig. 3, with Fig. 3a illustrating the cluster centres and Fig. reffig4 showing the evaluation metric scores.

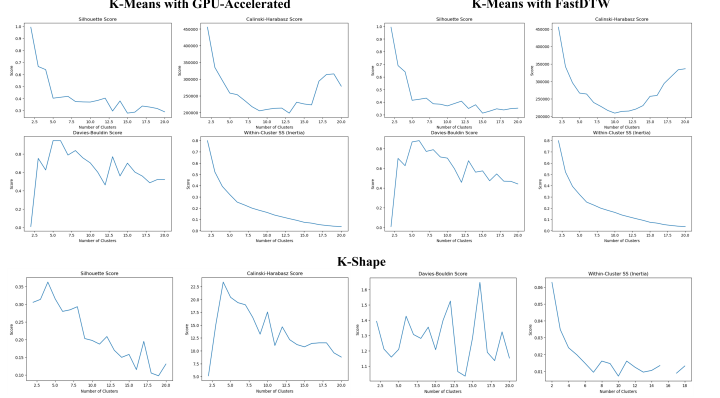
Fig. 3 reveals similarities in cluster centre locations, despite differences in cluster shapes. Fig. 3b mirrors the fluctuations across metrics, with GPU k-means and K-Shape exhibiting relatively superior performance on this dataset according to the index values. The Dunn index is excluded from Fig. 3b due to consistently zero values.

DBSCAN is another seminal density-based clustering algorithm that weakens in high-dimensional multivariate data, and factors like non-uniform density and varying cluster spacing can degrade performance. Analogous to k-means, DBSCAN is paired with Euclidean and FastDTW distances for clustering the stock dataset.

Given the large data volume, GPU acceleration and parallelization are implemented for efficiency. Four internal validation metrics are utilized for evaluation - Silhouette Score, Calinski-Harabasz Index, Davies-Bouldin Index, and Mean Distance to Nearest Cluster Member. Inertia is excluded as



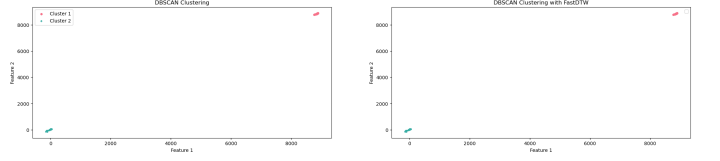
(a) K-means with GPU Acceleration, K-means with FastDTW, and K-shape Clustering Results



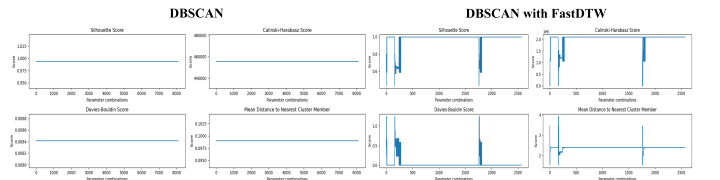
(b) K-means with GPU Acceleration, K-means with FastDTW Evaluation Metrics Result

Fig. 3: K-means with GPU Acceleration, K-means with FastDTW, and K-shape

DBSCAN does not explicitly minimize squared within-cluster distances.



(a) DBSCAN and DBSCAN with FastDTW Clustering Results



(b) DBSCAN and DBSCAN with FastDTW Evaluation Metrics Result

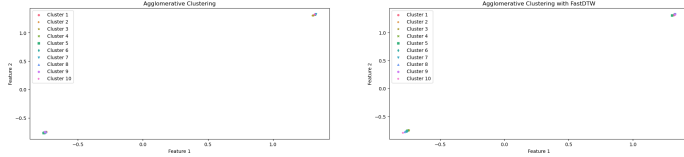
Fig. 4: DBSCAN and DBSCAN with FastDTW

Fig. 4 shows that DBSCAN with FastDTW significantly outperforms the Euclidean. This highlights the benefits of time-shaped distances for temporal data.

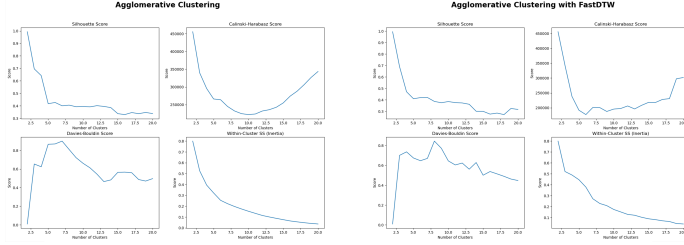
Agglomerative Clustering, a classic hierarchical model, faces limitations due to high computational complexity and singular value influence. As with DBSCAN, it combines with Euclidean and FastDTW to analyze the data, employing corresponding evaluation indices and reports.

A precomputed FastDTW distance matrix is integrated into Agglomerative Clustering to mitigate computational complex-

ity. Results for both models are presented in Figure 5a, while evaluation metrics appear in Figure 5b. Figure 5 suggests that the FastDTW is better for this dataset when combined with Agglomerative Clustering.



(a) Agglomerative Clustering and Agglomerative Clustering with Fast-DTW Clustering Results



(b) Agglomerative Clustering and Agglomerative Clustering with Fast-DTW Evaluation Metrics Result

Fig. 5: Agglomerative Clustering and Agglomerative Clustering with FastDTW

The proposed Voformer-EC neural network model proposes to achieve superior multivariate time series clustering performance while maintaining efficiency. Voformer-EC is evaluated using the same internal validation metrics applied to the traditional clustering benchmarks.

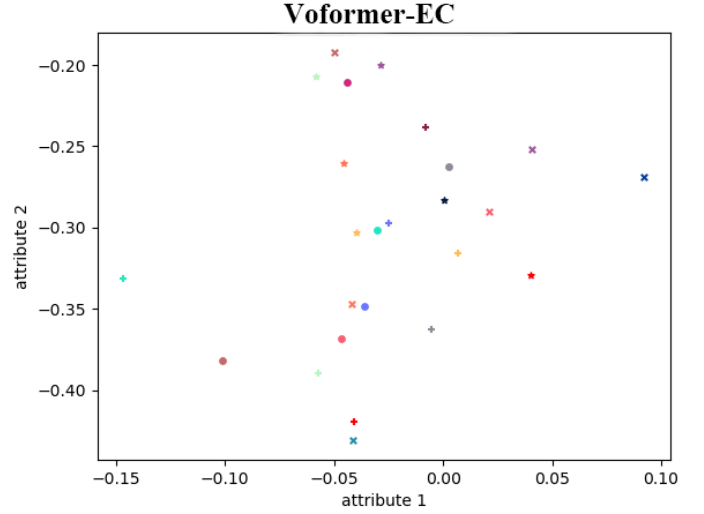
The evaluation index results are shown in Fig. 6. Voformer-EC demonstrates faster convergence and more stable clustering performance compared to traditional methods. Notably, the non-zero Dunn index indicates Voformer-EC identifies more clearly separated, compact clusters by effectively capturing complex relationships in time series data.

Benchmarking verifies Voformer-EC’s state-of-the-art clustering quality and efficiency for multivariate stock index data. The specialized neural architecture provides significant enhancements over traditional clustering techniques on key indicators of clustering performance. Leveraging the strengths of both neural networks and Extreme Clustering, Voformer-EC demonstrates rapid convergence and stable clustering performance on the time series feature matrix after model training.

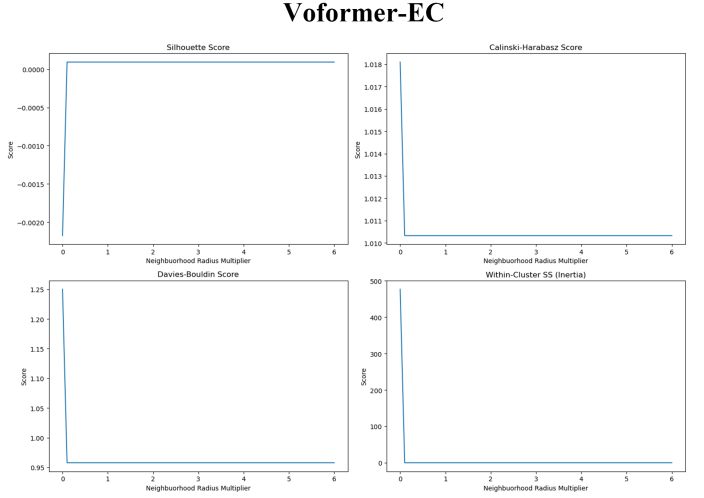
IV. RESULTS AND EVALUATION

This research aimed to propose an effective solution for multivariate time series clustering with application to stock market data. The proposed Voformer-EC framework integrates Informer-based neural networks, a volatility-aware activation function, and the Extreme Clustering algorithm.

Rigorous benchmarking was undertaken against traditional clustering techniques, including k-means, DBSCAN, hierarchical clustering, and variants incorporating Euclidean, DTW, and K-Shape distances. Multiple internal validation metrics quantified clustering performance.



(a) Voformer-EC Clustering Result



(b) Voformer-EC Evaluation Metrics Result

Fig. 6: Voformer-EC Clustering Result and Indexes Display

Experiments on a multivariate stock index dataset demonstrated Voformer-EC’s superior convergence, stability, and cluster separation compared to the baselines. Since Voformer-EC has the characteristics of both neural network and Extreme Clustering, it has fast convergence and very stable characteristics after convergence when performing cluster analysis on the extracted time feature matrix after training. Qualitative analysis of the clustered price trends also highlighted interpretable groupings of similarly behaving assets over the one-month period. This demonstrates Voformer-EC’s practical applicability for financial data analytics.

Limitations of the study include the lack of ground truth labels and a comparatively small single dataset for evaluation. Additional real-world datasets could help strengthen the empirical results. Hyperparameter tuning and ablation studies would also be beneficial. Next, we will test on more data sets to test more suitable datasets and customize more functions

for the Voformer-EC model on different specific datasets so that we can observe the results more intuitively and evaluate performance.

This research makes notable contributions through the novel Voformer-EC framework, rigorous benchmarking, and demonstrative experiments. It provides a robust, efficient multivariate time series clustering solution outperforms established techniques. The model shows promise for diverse temporal applications across climate, healthcare, finance, and more.

DECLARATION OF COMPETING INTEREST

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

ACKNOWLEDGEMENT

We sincerely thank the Climatic Data Centre, part of the National Meteorological Information Centre (CMA Meteorological Data Centre), for their invaluable assistance and cooperation in providing us with the meteorological data used in this study.

REFERENCES

- [1] Javed, A., Lee, B.S., Rizzo, D.M., 2020. A benchmark study on time series clustering. *Machine Learning with Applications* 1, 100001. URL: <https://www.sciencedirect.com/science/article/pii/S2666827020300013>, doi:10.1016/j.mlwa.2020.100001.
- [2] Alqahtani, A., Ali, M., Xie, X., Jones, M.W., 2021b. Deep time series clustering: A review. *Electronics* 10, 3001. URL: <https://www.mdpi.com/2079-9292/10/23/3001>, doi:10.3390/electronics10233001. number: 23 Publisher: Multidisciplinary Digital Publishing Institute.
- [3] Aghabozorgi, S., Seyed Shirkhorshidi, A., Ying Wah, T., 2015. Time-series clustering – a decade review. *Information Systems* 53, 16–38. URL: <https://www.sciencedirect.com/science/article/pii/S0306437915000733>, doi:10.1016/j.is.2015.04.007.
- [4] Levantesi, S., Nigri, A., Piscopo, G., 2022. Clustering-based simultaneous forecasting of life expectancy time series through long-short term memory neural networks. *International Journal of Approximate Reasoning* 140, 282–297. URL: <https://www.sciencedirect.com/science/article/pii/S0888613X21001730>, doi:10.1016/j.ijar.2021.10.008.
- [5] Bandara, K., Bergmeir, C., Smyl, S., 2020. Forecasting across time series databases using recurrent neural networks on groups of similar series: A clustering approach. *Expert Systems with Applications* 140, 112896. URL: <https://www.sciencedirect.com/science/article/pii/S0957417419306128>, doi:10.1016/j.eswa.2019.112896.
- [6] D’Urso, P., De Giovanni, L., Massari, R., 2021. Trimmed fuzzy clustering of financial time series based on dynamic time warping. *ANNALS OF OPERATIONS RESEARCH* 299, 1379–1395. doi:10.1007/s10479-019-03284-1.
- [7] Punyapornwithaya, V., Jampachaisri, K., Klaharn, K., Sansamur, C., 2021. Forecasting of milk production in northern thailand using seasonal autoregressive integrated moving average, error trend seasonality, and hybrid models. *FRONTIERS IN VETERINARY SCIENCE* 8, doi:10.3389/fvets.2021.775114.
- [8] Wang, M., Pan, J., Li, X., Li, M., Liu, Z., Zhao, Q., Luo, L., Chen, H., Chen, S., Jiang, F., Zhang, L., Wang, W., Wang, Y., 2022. Arima and arima-ernn models for prediction of pertussis incidence in mainland china from 2004 to 2021. *BMC PUBLIC HEALTH* 22, doi:10.1186/s12889-022-13872-9.
- [9] Song, Y., Cao, J., 2022. An arima-based study of bibliometric index prediction. *ASLIB JOURNAL OF INFORMATION MANAGEMENT* 74, 94–109. doi:10.1108/AJIM-03-2021-0072.
- [10] Moskolai, W.R., Abdou, W., Dipanda, A., Kolyang, 2021. Application of deep learning architectures for satellite image time series prediction: A review. *REMOTE SENSING* 13, doi:10.3390/rs13234822.
- [11] Belhadi, A., Djenouri, Y., Norvag, K., Ramampiaro, H., Masseglia, F., Lin, J.C.W., 2020. Space-time series clustering: Algorithms, taxonomy, and case study on urban smart cities. *ENGINEERING APPLICATIONS OF ARTIFICIAL INTELLIGENCE* 95, doi:10.1016/j.engappai.2020.103857.
- [12] Liu, H., Liu, Y., Zhang, R., Wu, X., 2021a. A clustering algorithm via density perception and hierarchical aggregation based on urban multimodal big data for identifying and analyzing categories of poverty-stricken households in china. *Scientific Programming* 2021, 13. URL: <https://doi.org/10.1155/2021/6692975>, doi:10.1155/2021/6692975.
- [13] Cardarilli, G.C., Di Nunzio, L., Fazzolari, R., Nannarelli, A., Re, M., Spano, S., 2020. N-dimensional approximation of euclidean distance. *IEEE TRANSACTIONS ON CIRCUITS AND SYSTEMS II-EXPRESS BRIEFS* 67, 565–569. doi:10.1109/TCSII.2019.2919545.
- [14] Li, Y., Shen, D., Nie, T., Kou, Y., 2022a. A new shape-based clustering algorithm for time series. *INFORMATION SCIENCES* 609, 411–428. doi:10.1016/j.ins.2022.07.105.
- [15] Herrmann, M., Webb, G.I., 2023. Amercing: An intuitive and effective constraint for dynamic time warping. *PATTERN RECOGNITION* 137, doi:10.1016/j.patcog.2023.109333.
- [16] Ge, L., Chen, S., 2020. Exact dynamic time warping calculation for weak sparse time series. *APPLIED SOFT COMPUTING* 96, doi:10.1016/j.asoc.2020.106631.
- [17] Zhou, H., Zhang, S., Peng, J., Zhang, S., Li, J., Xiong, H., Zhang, W., 2020. Informer: Beyond efficient transformer for long sequence time-series forecasting. URL: <https://arxiv.org/abs/2012.07436>, doi:10.48550/ARXIV.2012.07436.
- [18] Zhou, H., Zhang, S., Peng, J., Zhang, S., Li, J., Xiong, H., Zhang, W., 2020. Informer: Beyond efficient transformer for long sequence time-series forecasting. URL: <https://arxiv.org/abs/2012.07436>, doi:10.48550/ARXIV.2012.07436.
- [19] Wu, Z., Pan, F., Li, D., He, H., Zhang, T., Yang, S., 2022. Prediction of photovoltaic power by the informer model based on convolutional neural network. *SUSTAINABILITY* 14, doi:10.3390/su142013022.
- [20] Kayim, F., Yilmaz, A., 2022b. Time series forecasting with volatility activation function. *IEEE Access* 10, 104000–104010. doi:10.1109/ACCESS.2022.3211312. Lederer, J., 2021. Activation functions in artificial neural networks: A systematic overview. *arXiv preprint arXiv:2101.09957*.
- [21] Wang, S., Li, Q., Zhao, C., Zhu, X., Yuan, H., Dai, T., 2021b. Extreme clustering – A clustering method via density extreme points. *Information Sciences* 542, 24–39. URL: <https://www.sciencedirect.com/science/article/pii/S0020025520306587>, doi:10.1016/j.ins.2020.06.069.
- [22] Hou, J., Zhang, A., Qi, N., 2020. Density peak clustering based on relative density relationship. *PATTERN RECOGNITION* 108, doi:10.1016/j.patcog.2020.107554.
- [23] NSE - Nifty 50 Index Minute data (2015 to 2022), 2023. Kaggle CC0: Public Domain. Kaggle. URL: <https://www.kaggle.com/datasets/debashis74017/nifty-50-minute-data>.
- [24] Ezugwu, A.E., Ikotun, A.M., Oyelade, O.O., Abualigah, L., Agushaka, J.O., Eke, C.I., Akinyelu, A.A., 2022. A comprehensive survey of clustering algorithms: State-of-the-art machine learning applications, taxonomy, challenges, and future research prospects. *Engineering Applications of Artificial Intelligence* 110, 104743. URL: <https://www.sciencedirect.com/science/article/pii/S095219762200046X>, doi:10.1016/j.engappai.2022.104743.