

Localization of Global Forecasting Models with a Clustering Approach

Hossein Abbasimehr¹  and Mohammad Khodizadeh-Nahari² 

1. Associate Professor, Faculty of Information Technology and Computer Engineering, Azarbaijan Shahid Madani University, Tabriz, Iran, Email: abbasimehr@azaruniv.ac.ir
2. Corresponding Author, Assistant Professor, Faculty of Information Technology and Computer Engineering, Azarbaijan Shahid Madani University, Tabriz, Iran, Email: m.khodizadeh@azaruniv.ac.ir

Article Info	ABSTRACT
Article type: Research Article Article history: Received 5 November 2024 Received in revised form 10 February 2025 Accepted 4 August 2025 Published online 21 August 2025 Keywords: Time series forecasting, Global forecasting model, Time series clustering, Long short term memory network.	With the increasing generation of time series data, forecasting models trained on a set of time series, known as global forecasting models, outperform univariate forecasting models trained on individual series. However, the performance of global models may decrease when faced with heterogeneous data sets of time series with different lengths. In this study, a new method for localization of clustering-based global forecasting models is presented. The main steps of the proposed method include (1) extracting relevant features from each time series (2) clustering time series based on features extracted using K-Medoids and spectral clustering algorithms (3) implementing a global forecasting model using a combination of Temporal Convolution Network and its training for each cluster. To evaluate the prediction accuracy of the proposed approach, experiments were conducted on the M3 dataset that contains 1426 time series with unequal-length. The results of the experiments show the superior performance of the proposed clustering-based models compared to the baseline models and the benchmark models. The proposed model has 0.57 less error in terms of SMAPE metric.
Cite this article: Abbasimehr, H. & Khodizadeh-Nahari, M. (2025)., Localization of Global Forecasting Models with a Clustering Approach. <i>Enginiiring Management and Soft Computing</i>, 11 (1). 134-155. DOI: https://doi.org/10.22091/jemsc.2025.11595.1218	
	© The Author(s) DOI: 10.22091/jemsc.2025.11595.1218 Publisher: University of Qom

1) INTRODUCTION

A global forecasting model is a model that is obtained by considering samples from several time series. Compared to separate models for each series, these models perform better in predicting related time series. As a result of this evolution, the tendency to use deep learning methods as predictive models of time series has become more popular than other statistical methods. Considering that most time series are not homogeneous, the model obtained using the whole series usually does not have good accuracy. For this reason, localization of global models is a method to increase the accuracy of global forecasting models. The motivation behind this study is to propose a new method for the localization of global forecast models based on clustering.

2) Solution Method

In this research, we use a clustering-based approach to localize global prediction models. The proposed method generally includes three stages: pre-processing, global forecasting, and post-processing. The pre-processing step comprises extracting the features of the time series for clustering, clustering the time series based on the extracted features. The global forecasting stage includes special pre-processing operations, such as mean normalization, logarithmic transformation, and de-seasonalization, as well as training a forecasting model based on the TCN-CNN model. The topology of the model is represented in Table 1. In the post-processing stage, the obtained forecasts are returned to their original scale by considering the parameters calculated in the pre-processing. In this research, we use the M3 dataset, which includes 1426 time series, to evaluate the proposed method. For clustering time series, we employed K-Medoids (CL.KM) and Spectral clustering (CL.SP) algorithms. The results of experiments indicated that clustering-based models have better results than their baseline counterparts in sMAPE and MASE error measures. Also, the results show the superiority of the proposed models over some benchmarks.

Table1. Topology of TCN-CNN

Layers
Input (shape: (1, Lag (input window)))
TCN (filters: 64, kernel size: 3, dilations: (1, 2, 4, 8))
1D Convolutional (filters: 64, kernel size: 4, activation: linear)
1D Convolutional (filters: 16, kernel size: 4, activation: linear)
Flatten
Dense (units: Forecast horizon, activation: linear)
Dense

3) Equations

We use relations 1- 4 for normalization, logarithmic transformation, de-seasonalization.

$$X_{i,normalized} = \frac{X_i}{\frac{1}{k} \sum_{t=1}^k X_{i,k}} \quad (1)$$

k: number of points in time series X_i .

$$T_{i,normalized} \& \log_trans = \begin{cases} \log(T_{i,normalized}), & \min(TS) > 0; \\ \log(T_{i,normalized} + 1), & \min(TS) = 0, \end{cases} \quad (2)$$

$$T_{i,normalized} \& \log_trans = s_i + t_i + r_i \quad (3)$$

$$T_{i,normalized} \& \log_{trans} \& deseasonalized = T_{i,normalized} \& \log_trans - s_i \quad (4)$$

s_i : seasonal components, t_i : trend components and r_i residual components

For model evaluation we used two major metric: sMAPE (relation 5) and MASE (relation 6)

$$sMAPE = \frac{200}{h} \sum_{i=1}^h \left(\frac{|A_i - F_i|}{|A_i| + |F_i|} \right) \quad (5)$$

$$\text{MASE} = \frac{1}{h} \frac{\sum_{i=1}^h (|A_i - F_i|)}{\frac{1}{n-M} \sum_{i=s+1}^n (|A_i - A_{i-s}|)} \quad (6)$$

A_i: Actual value

F_i: Forecasted value

4) Conclusion

The results of this study indicated that clustering-based global models (CL, KM and CL.SP models in Table 2) that train a separate model per cluster outperform their baseline counterparts, in which a single model is trained using all time series data. Time series clustering identifies homogenous clusters, which leads to accuracy improvement. Also, we investigated the effect of de-seasonalization and showed that it positively impacts the performance of both clustering-based models and baseline ones.

Table 2. Comparison of the proposed model with competing models

MASE	sMAPE	Model
0.894	13.94	CL.KM
0.898	13.99	CL.SP
0.914	14.51	FFORMA
1.065	16.41	CatBoost
1.167	15.74	DeepAR
0.934	14.76	N-BEATS

REFERENCES

- Bandara, K., C. Bergmeir, and S. Smyl, Forecasting across time series databases using recurrent neural networks on groups of similar series: A clustering approach. *Expert Systems with Applications*, 140: p. 112896, (2020), <https://doi.org/10.1016/j.eswa.2019.112896>
- Godahewa, R., et al., Ensembles of localised models for time series forecasting. *Knowledge-Based Systems*, 233: p. 107518, (2021), <https://doi.org/10.1016/j.knosys.2021.107518>
- Bandara, K., Forecasting with Big Data Using Global Forecasting Models, in *Forecasting with Artificial Intelligence: Theory and Applications*, M. Hamoudia, S. Makridakis, and E. Spiliotis, Editors. Springer Nature Switzerland: Cham. p. 107-122, (2023), <https://doi.org/10.1007/978-3-031-35879-1>.

محل سازی مدل های پیش بینی سراسری با یک رویکرد خوشه بندی

حسین عباسی مهر^۱ و محمد خودی زاده نهاری^۲  

۱. دانشیار، دانشکده فناوری اطلاعات و مهندسی کامپیوتر، دانشگاه شهید مدنی آذربایجان، تبریز، ایران، رایانامه: abbasimehr@azaruniv.ac.ir

۲. نویسنده مسئول، استادیار، دانشکده فناوری اطلاعات و مهندسی کامپیوتر، دانشگاه شهید مدنی آذربایجان، تبریز، ایران، رایانامه: m.khodizadeh@azaruniv.ac.ir

اطلاعات مقاله	چکیده
نوع مقاله: مقاله پژوهشی تاریخ دریافت: ۱۴۰۳/۰۸/۱۵ تاریخ بازنگری: ۱۴۰۳/۱۱/۲۲ تاریخ پذیرش: ۱۴۰۴/۰۵/۱۳ تاریخ انتشار: ۱۴۰۴/۰۵/۳۰ کلیدواژه ها: پیش بینی سری زمانی، خوشه بندی سری زمانی، شبکه حافظه کوتاه مدت طولانی، مدل پیش بینی سراسری	با تولید روزافزون داده های سری زمانی، مدل های پیش بینی که بر روی مجموعه ای از سری های زمانی آموزش داده می شوند و به عنوان مدل های پیش بینی سراسری شناخته می شوند، نسبت به مدل های پیش بینی تک متغیره که بر روی سری های جداگانه آموزش می بینند، عملکرد بهتری دارند. با این حال، عملکرد مدل های سراسری ممکن است در مواجهه با مجموعه داده های ناهمگن سری های زمانی با طول های متفاوت، کاهش یابد. در این مطالعه، یک روش جدید برای محل سازی مدل های پیش بینی سراسری مبتنی بر خوشه بندی ارائه می شود. مراحل اصلی روش ارائه شده شامل: (۱) استخراج ویژگی های مرتبط از هر سری زمانی، (۲) خوشه بندی سری های زمانی بر پایه ویژگی های استخراج شده با استفاده از الگوریتم های K-Medoids و خوشه بندی طیفی، (۳) پیاده سازی دو مدل پیش بینی سراسری با استفاده از مدل مبتنی بر شبکه کانولوشنی زمانی و مدل مبتنی بر مکانیزم توجه چندگانه می باشد. برای ارزیابی دقت پیش بینی، آزمایشاتی بر روی مجموعه داده M3 شامل ۱۴۲۸ سری زمانی با طول های غیریکسان انجام شد. نتایج آزمایشات، نشان دهنده عملکرد برتر مدل های پیشنهادی در مقایسه با مدل های پایه و مدل های محک است. در معیار SMAPE، مدل پیشنهادی ۰.۵۷ خطای کمتری دارد.

استناد: عباسی مهر، حسین و محمد، خودی زاده نهاری. (۱۴۰۴). «محل سازی مدل های پیش بینی سراسری با یک رویکرد خوشه بندی». مدیریت مهندسی

و رایانش نرم، دوره ۱۱ (۱). صص: ۱۵۵-۱۳۴. <https://doi.org/10.22091/jemsc.2025.11595.1218>



© نویسندگان.

ناشر: دانشگاه قم

(۱) مقدمه

با توجه به رشد روزافزون تولید داده در صنایع مختلف از جمله انرژی، خرده‌فروشی و گردشگری، مدل‌های پیش‌بینی سری‌های زمانی نقش بسیار حیاتی و حساسی را در این صنایع ایفا می‌کنند (گوداهیوا^۱ و همکاران، ۲۰۲۱). به همین دلیل و با توجه به پیچیدگی روزافزون مسائل پیش‌بینی سری‌های زمانی، نیاز است این مدل‌ها از پیش‌بینی سری‌های زمانی به صورت جداگانه به پیش‌بینی مجموعه‌ای از سری‌های زمانی مرتبط به هم تکامل یابند (باندرا^۲ و همکاران، ۲۰۲۰). در صنایع مختلف، انجام پیش‌بینی‌های دقیق نقش بسیار حیاتی در تصمیم‌گیری‌های تجاری ایفا می‌کند. برای مثال: در صنعت خرده‌فروشی مانند فروشگاه‌های بزرگی چون آمازون^۳ و یا الومارت^۴، پیش‌بینی میزان تقاضا برای تنظیم صحیح موجودی کالاها، نقش بسزایی دارد. در حوزه تامین انرژی، پیش‌بینی میزان مصرف در میزان تخصیص منابع ارزشمند سوختی و مدیریت بازار انرژی، بسیار حیاتی است. همچنین در صنعت گردشگری و خدمات حوزه سلامت، پیش‌بینی میزان تقاضا برای خدمات در مناطق جغرافیایی مختلف از اهمیت بالایی برخوردار است (هیوآمالاگ^۵ و همکاران، ۲۰۲۲ و مارتینز^۶). و همکاران، ۲۰۱۸) تاکنون مطالعات بسیاری در پیش‌بینی سری‌های زمانی انجام شده است. در این زمینه، اکثر تحقیقات بر پیش‌بینی سری‌های زمانی تکی متمرکز بوده‌اند. بطور کلی این تحقیقات دو دسته از روش‌ها را مورد استفاده قرار داده‌اند. دسته اول شامل: روش‌های سنتی آماری نظیر میانگین متحرک^۷، هموارسازی نمایی^۸ و روش ARIMA^۹ (هیندمن^{۱۰} و همکاران، ۲۰۰۲ و باکس و همکاران ۲۰۱۵) که مهمترین مشکلی که روش‌های سنتی دارند، این است که برای پیش‌بینی سری‌های زمانی حاوی الگوهای خطی مناسب هستند و دسته دوم که برای مدل‌سازی الگوها و روابط غیرخطی موجود در داده‌ها مناسب هستند، اخیراً در روش‌های یادگیری ماشین و یادگیری عمیق کاربردهای موفق داشته‌اند (پارمیزان^{۱۱} و همکاران، ۲۰۱۹). در دسته روش‌های یادگیری ماشین، شبکه‌های عصبی مصنوعی (ANN)، ماشین‌های بردار پشتیبان (SVM) و K نزدیک‌ترین همسایه (KNNs) جزء پرکاربردترین روش‌ها هستند. همچنین گونه‌های مختلف روش‌های شبکه عصبی بازگشتی (RNN) نظیر شبکه حافظه کوتاه‌مدت طولانی (LSTM)، واحد گیتی بازگشتی (GRU)، عملکرد موفق در پیش‌بینی سری زمانی از خود نشان داده‌اند (هیوآمالاگ و همکاران، ۲۰۲۱). شبکه‌های عصبی بازگشتی، توانایی پردازش و مدل‌سازی وابستگی بین داده‌ها در طول زمان را دارند. اغلب مطالعات پیشین روی ساخت مدل بر پایه یک سری زمانی تکی متمرکز بوده‌اند (باندرا و همکاران، ۲۰۲۰). بزرگترین چالش استفاده از روش‌های یادگیری ماشین و یادگیری عمیق با اتخاذ این رویکرد، کافی نبودن نمونه‌های داده جهت آموزش مدل‌ها است.

یک مدل پیش‌بینی سراسری، مدلی است که با در نظر گرفتن نمونه‌هایی از چندین سری زمانی به دست می‌آید. این

¹ Godahewa

² Bandara

³ www.amazon.com

⁴ www.walmart.com

⁵ Hewamalage

⁶ Martínez

⁷ Moving average

⁸ Exponential Smoothing

⁹ Autoregressive integrated moving average

¹⁰ Hyndman

¹¹ Parmezan

مدل ها، در مقایسه با مدل های جداگانه، عملکرد بهتری در پیش بینی سری های زمانی مرتبط به هم دارند (گوداهیا و همکاران، ۲۰۲۱، هیوامالاگ و همکاران، ۲۰۲۲ و هیوامالاگ و همکاران، ۲۰۲۱). در نتیجه این تکامل، تمایل به استفاده از روش های یادگیری عمیق به عنوان مدل های پیش بینی کننده سری های زمانی نسبت به سایر روش های آماری، محبوبیت بیشتری یافته است. با توجه به اینکه در یک مجموعه سری زمانی، اغلب سری های زمانی همگن نیستند، معمولاً مدل به دست آمده با استفاده از کل سری ها، از دقت خوبی برخوردار نیست. به همین دلیل، محلی سازی مدل های سراسری یک روش برای افزایش دقت مدل های پیش بینی سراسری می باشد (گوداهیا و همکاران، ۲۰۲۱ و باندرا و همکاران، ۲۰۲۰).

در این مطالعه، یک روش جدید برای محلی سازی مدل های پیش بینی سراسری مبتنی بر خوشه بندی ارائه می شود. با استفاده از خوشه بندی سری های زمانی و افراز آنها به خوشه هایی شامل سری های زمانی مشابه، قادر خواهیم بود راه حلی برای مشکل محلی سازی مدل های پیش بینی سراسری ارائه دهیم. زیرا با استفاده از این روش می توان یک مدل سراسری پیش بینی برای هر خوشه آموزش داد که بتواند جزئیات دقیق تری از سری های زمانی مشابه را فرا گیرد.

مراحل اصلی روش ارائه شده شامل: (۱) استخراج ویژگی های مفید و مرتبط از هر سری زمانی، (۲) خوشه بندی سری های زمانی بر پایه ویژگی های استخراج شده با استفاده از الگوریتم های K-Medoids (شوبرت و روزیو^{۱۲}، ۲۰۲۱) و خوشه بندی طیفی^{۱۳} (پی و ساکورای^{۱۴}، ۲۰۱۶)، (۳) پیاده سازی دو مدل پیش بینی سراسری با ترکیب شبکه پیچشی زمانی^{۱۵} و شبکه پیچشی^{۱۶} و مدل مبتنی بر مکانیزم توجه چند گانه^{۱۷} است.

با توجه به خاصیت ناهمگونی موجود در مجموعه داده های سری زمانی، پیش پردازش سری ها و همچنین پس پردازش نتایج پیش بینی از اهمیت فوق العاده ای برخوردار است. در روش پیشنهادی در مرحله پس پردازش نتایج پیش بینی از مدل هموار سازی نمایی^{۱۸} جهت تخمین مولفه های فصلی استفاده می شود.

روش پیشنهادی روی مجموعه داده M3 پیاده سازی شد. نتایج آزمایشات، نشان دهنده برتری مدل های پیش بینی حاصل از رویکرد پیشنهادی می باشد. روش های مبتنی بر خوشه بندی نسبت به روش پایه دارای عملکرد پیش بینی بهتری هستند. همچنین مقایسه دقت مدل ها نسبت به مدل های موجود در تحقیقات فعلی حاکی از برتری مدل پیشنهادی است. بطور خلاصه نوآوری های این تحقیق به شرح زیر است:

۱. ارائه رویکردی جهت محلی سازی مدل های پیش بینی سراسری به منظور افزایش دقت پیش بینی در مجموعه داده های ناهمگن سری زمانی

۲. ارائه یک روش خوشه بندی مبتنی بر ویژگی جهت خوشه بندی سری های زمانی با طول متغیر

۳. پیاده سازی دو مدل پیش بینی مبتنی بر شبکه های عصبی عمیق

۴. ارزیابی رویکرد پیشنهادی با استفاده از ۱۴۲۸ سری زمانی مجموعه داده M3

¹² Schubert & Rousseeuw

¹³ Spectral clustering

¹⁴ Ye & Sakurai

¹⁵ Temporal Convolutional Networks (TCN)

¹⁶ Convolutional neural network

¹⁷ Multi-Head Attention

¹⁸ Exponential Smoothing

ساختار ادامه این مقاله بدین شرح است: در بخش ۲ به مرور ادبیات پیش‌بینی سراسری پرداخته می‌شود، در بخش ۳ روش پیشنهادی با جزئیات کامل توصیف خواهد شد. بخش ۴ شامل توصیف مجموعه داده، تنظیمات آزمایشات و تشریح معیارهای ارزیابی عملکرد، نتایج پیش‌بینی و تحلیل عملکرد مدل‌ها خواهد بود. در بخش ۵ نتیجه‌گیری بیان خواهد شد.

(۲) پیشینه تحقیق

در سال‌های اخیر، مطالعات بسیاری در حوزه پیش‌بینی سری‌های زمانی مبتنی بر یادگیری ماشین و مخصوصاً یادگیری عمیق صورت گرفته‌است (هیوآمالاگ و همکاران، ۲۰۲۱). بطور کلی می‌توان مدل‌های پیش‌بینی سری زمانی مبتنی بر رویکرد یادگیری ماشین را به دو دسته اصلی سراسری و محلی تقسیم کرد (گوداهیا و همکاران، ۲۰۲۱) که در میان آنها، مدل‌های پیش‌بینی سراسری (جانوفسکی^{۱۹} و همکاران ۲۰۲۰)، یک رویکرد نسبتاً جدید هستند و پس از استفاده به عنوان پایه و اساس حل مسئله در رقابت‌های M4 بسیار محبوب شدند (اسمیل^{۲۰}، ۲۰۲۰). برخلاف رویکردهای محلی مانند هموارسازی نمایی (ETS) (هیندمن و همکاران، ۲۰۰۸) و مدل‌های خود رگرسیون میانگین متحرک انباشته (ARIMA) (باکس^{۲۱} و همکاران، ۲۰۱۵) که در آنها یک مدل پیش‌بینی به ازای هر سری زمانی ساخته می‌شود، در رویکرد مدل‌های پیش‌بینی سراسری^{۲۲} یک مدل واحد با استفاده از تمام سری‌های زمانی ساخته می‌شوند (جانوفسکی و همکاران ۲۰۲۰). این مدل‌ها پتانسیل بسیار بالایی در ایجاد دقت بالای پیش‌بینی از خود نشان داده‌اند. برای مثال: در مسابقات معتبر M4 و M5 (ماکریداکیس^{۲۳} و همکاران، ۲۰۲۲ و ماکریداکیس و همکاران، ۲۰۱۸) با ارائه دقت پیش‌بینی بالا، این مدل‌ها توانستند برنده مسابقات شوند. با داشتن یک مجموعه داده با p سری زمانی $X = \{X_1, X_2, \dots, X_p\}$ ، در پیش‌بینی سراسری هدف، ساخت مدل F با در نظر گرفتن تمامی داده‌های مجموعه داده است. مدل سراسری F به این صورت تعریف می‌شود:

$$X_i^M = F(X, \theta) \quad (۱)$$

که در آن M نشان‌دهنده افق پیش‌بینی است. i نشان‌دهنده سری زمانی موردنظر است که مقادیر آینده آن توسط مدل F پیش‌بینی خواهد شد و θ پارامترهای مدل F است که با اطلاعات تمامی سری‌های زمانی تعیین می‌شود. برای مثال: در یک شبکه عصبی، θ نشان‌دهنده وزن‌ها و بایاس است. بنابراین براساس مدل پایه‌ای که در ساخت مدل پیش‌بینی استفاده می‌شود، θ تغییر می‌کند. در صورت استفاده از رویکرد سنتی، پیش‌بینی سری زمانی شامل: ساخت مدل‌های جداگانه به ازای هر سری خواهد شد. در نتیجه دارای p مدل $F = \{F_1, F_2, \dots, F_p\}$ خواهیم بود که پارامترهای هر مدل با داده‌های هر سری به دست می‌آیند.

مدل‌های سراسری این قابلیت را دارند که در مجموعه‌ای از داده‌های سری زمانی با طول‌های متفاوت اعمال شوند. در مواردی که سری‌های زمانی طول‌های متفاوتی داشته باشند، تعداد نمونه‌های تولیدشده برای هر سری زمانی ممکن است

¹⁹ Januschowski

²⁰ Smyl

²¹ Box

²² Global Forecasting Models (GFM)

²³ Makridakis

متفاوت می شود. با این حال، چنین تناقضی، مانعی در روند آموزش مدل های سراسری به وجود نمی آورد. این یک تمایز قابل توجه بین GFM ها و مدل های پیش بینی تک متغیره است.

شبکه های عصبی بازگشتی (RNN) به دلیل توانایی طبیعی آنها در پردازش داده های متوالی برای پیش بینی سری های زمانی نتایج خوبی داشته اند. در حوزه مدل های سراسری، RNN ها جزء مدل های یادگیری عمیق پُر کاربرد می باشند (هیوامالاگ و همکاران، ۲۰۲۱). اسمیل و کوبر^{۲۴} (۲۰۱۶) معماری جدید چندلایه با شبکه LSTM معرفی کردند که شامل اتصالات پرش برای کاهش مشکل ناپدید شدن گرادیان است. رویکرد آنها همچنین دارای ویژگی های کمکی به دست آمده از مدل هموارسازی نمایی (ETS) بود که بطور قابل توجهی دقت مدل LSTM را افزایش داد. نتایج تجربی آنها نشان داد که مدل LSTM بهبود یافته با ETS، از مدل استاندارد LSTM بهتر عمل می کند. در ادامه به تشریح چندین مطالعه مهم دیگر در زمینه مدل های پیش بینی سراسری می پردازیم.

باندرا و همکاران (۲۰۲۰)، از یک معماری LSTM چندلایه^{۲۵} در کنار تکنیک های خوشه بندی سری های زمانی برای محلی سازی GFM ها و حل مشکل ناهمگونی در داده های سری زمانی استفاده کرده اند. روش پیشنهادی آنها، عملکرد بهتری نسبت به مدل LSTM که در آن یک مدل واحد روی تمام داده های سری زمانی آموزش داده می شود، از خود نشان داده است.

در مطالعه ای، سالیناس^{۲۶} و همکاران (۲۰۲۰) تکنیک DeepAR را ارائه کردند، یک مدل دنباله به دنباله^{۲۷} با معماری رمزگشای RNN که برای پیش بینی احتمالی طراحی شده است و قابلیت پیش بینی سری های زمانی مرتبط به هم را دارد. لاپتئو^{۲۸} و همکاران (۲۰۱۷) استفاده از شبکه رمزگذار خودکار برای آموزش یک مدل سراسری در سری های زمانی متنوع را پیشنهاد داده اند. رمزگذار خودکار ویژگی هایی را از سری های زمانی استخراج می کند که این ویژگی ها با ورودی خام برای پیش بینی نهایی ترکیب می شود.

اسمیل (۲۰۲۰) یک مدل ترکیبی مبتنی بر هموارسازی نمایی و شبکه RNN (ES-RNN) را معرفی کرده است. این مدل شامل پارامترهای محلی (خاص هر سری) و سراسری است که امکان یادگیری مدل سراسری با در نظر گرفتن ویژگی های منحصر به فرد هر سری را فراهم می کند. این روش برنده رقابت پیش بینی M4 شده است.

باندرا و همکاران (۲۰۲۱) افزایش دقت GFM را از طریق داده افزایی و یادگیری انتقالی پیشنهاد کردند. آنها از تکنیک های مختلف داده افزایی سری های زمانی برای تکمیل سری های زمانی استفاده کردند سپس از داده های تولید شده برای آموزش مدل از طریق رویکردهای یادگیری تلفیقی و انتقالی استفاده کردند. یافته های آنها نشان داد که این روش ها از مدل های سراسری پایه و چندین تکنیک پیش بینی تک متغیره پیشرفته به ویژه در مجموعه داده های کوچک تا متوسط، بهتر عمل می کنند.

گوداهوا و همکاران، (۲۰۲۱) تکنیک های مختلف یادگیری گروهی با مدل های مختلف پیشنهاد داده اند و نتایج،

²⁴ Smyl & Kuber

²⁵ Stacked

²⁶ Salinas

²⁷ Seq-to-Seq

²⁸ Laptev

حاکمی از برتری مدل‌های یادگیری گروهی نسبت به مدل‌های سنتی و یادگیری عمیق دارد. باندرا (۲۰۲۳) رویکردهای مختلفی برای پرداختن به رانش مفهومی^{۲۹} در پیش‌بینی سری‌های زمانی پیشنهاد داده‌است. چالشی که در مدل‌های پیش‌بینی سراسری وجود دارد این است که این مدل‌ها به اندازه کافی برای سری‌های زمانی خاص، محلی نیستند به‌ویژه زمانی که مجموعه سری‌های زمانی ناهمگن هستند. رویکردی کاربردی‌تر و عمومی‌تر برای پرداختن به این چالش که توانایی استفاده در هرگونه مدل پیش‌بینی سراسری را داشته باشد، می‌تواند توسعه مدل‌هایی برای زیرمجموعه‌هایی از تمام سری‌های زمانی یک مجموعه داده باشد. این رویکرد قادر است بطور همزمان پیچیدگی مدل را کنترل کرده و مدل‌های محلی‌تری بسازد که بر روی گروه‌های خاصی از سری‌های زمانی متمرکز می‌شوند. یک روش پیشنهادی برای اجرای این رویکرد، می‌تواند استفاده از خوشه‌بندی سری‌های زمانی هم به‌عنوان راهی برای ساختن مدل تجمیعی پیچیده کنترل‌شده و هم برای توجه به مسئله ناهمگونی باشد. مونتر و هیندمن^{۳۰} (۲۰۲۱) نشان داده‌اند که خوشه‌بندی تصادفی سری‌های زمانی می‌تواند با بالا بردن میزان پیچیدگی مدل، دقت پیش‌بینی را بهبود بخشد. برای بررسی ناهمگونی داده‌ها در (باندرا و همکاران، ۲۰۲۰) براساس روشی که پیشنهاد شده‌است، ابتدا سری‌های زمانی خوشه‌بندی می‌شوند سپس مدل‌های پیش‌بینی سراسری جداگانه‌ای برای هر خوشه به‌دست آمده، آموزش داده می‌شوند.

(۳) روش تحقیق

در این تحقیق از یک رویکرد مبتنی بر خوشه‌بندی برای محلی‌سازی مدل‌های پیش‌بینی سراسری استفاده می‌کنیم. شکل ۱ مراحل روش پیشنهادی را توصیف می‌کند. همانطور که در این شکل مشاهده می‌شود، روش پیشنهادی شامل سه مرحله پیش‌پردازش، ساخت مدل پیش‌بینی سراسری و پس‌پردازش است. مرحله پیش‌پردازش شامل استخراج ویژگی‌های زمانی برای خوشه‌بندی، خوشه‌بندی سری‌های زمانی مبتنی بر ویژگی‌های استخراج شده است. مرحله ساخت مدل پیش‌بینی شامل عملیات پیش‌پردازش خاص این مرحله مانند نرمال‌سازی با میانگین، تبدیل لگاریتمی، حذف مولفه فصلی است. در مرحله پس‌پردازش^{۳۱} مقادیر پیش‌بینی به‌دست آمده توسط مدل پیش‌بینی با در نظر گرفتن پارامترهای محاسبه شده در پیش‌پردازش به مقیاس اولیه خود بازگردانده می‌شوند. در ادامه جزئیات هر کدام از مراحل شرح داده می‌شود.

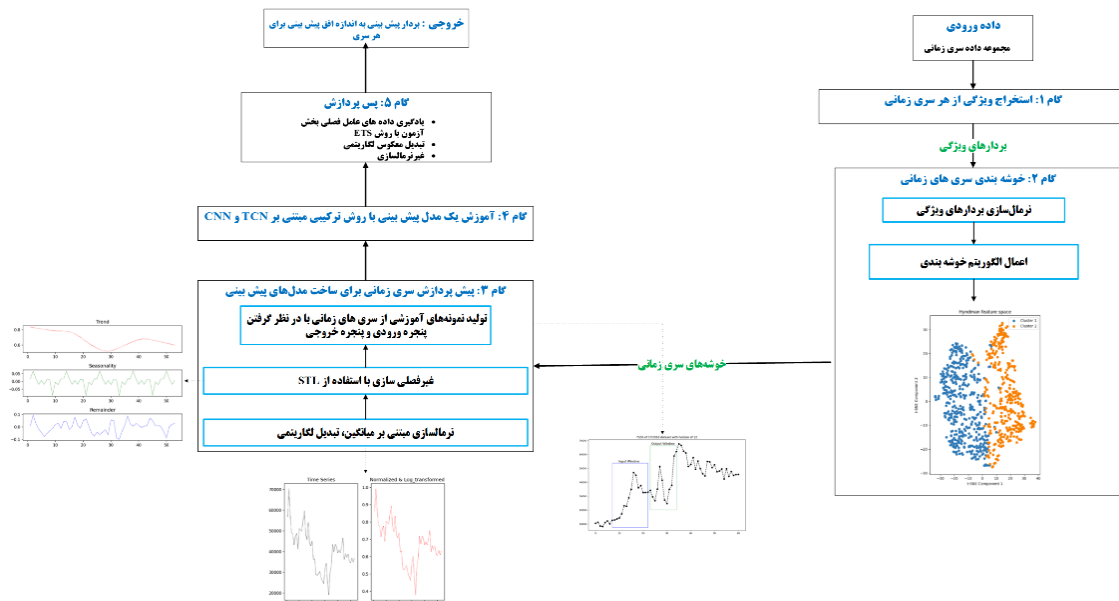
۳-۱) استخراج ویژگی‌های سری زمانی

با توجه به اینکه اغلب در مجموعه داده‌های سری زمانی، با سری‌های دارای طول‌های غیریکسان سروکار داریم، برای گروه‌بندی سری‌های شبیه به هم خوشه‌بندی مبتنی بر ویژگی را می‌توان در نظر گرفت. ویژگی‌های استخراج شده از هر سری زمانی به صورت یک بردار در نظر گرفته می‌شود. در این مطالعه از ویژگی معرفی شده توسط هیندمن و همکاران (۲۰۱۹) استفاده می‌گردد (جدول ۱).

²⁹ Concept drift

³⁰ Montero & Hyndman

³¹ Post processing



شکل ۱. رویکرد پیشنهادی برای محلی سازی مدل های پیش بینی سراسری

جدول ۱. ویژگی های مستخرج از سری های زمانی برای خوشه بندی

ویژگی	توضیحات
Mean	میانگین سری زمانی
Var	واریانس سری زمانی
ACF-1x	مرتبه اول همبستگی
Trend	شدت روند در سری زمانی
Linearity	شدت خطی بودن سری زمانی
Curvature	شدت انحنای سری زمانی
Entropy	شدت آنتروپی
Lumpiness	تغییر واریانس در باقیمانده
Spikiness	مقادیر غیرمنتظره سری زمانی
Lshift	تغییر سطح با استفاده از پنجره لغزان
Vchange	تغییر واریانس
Fspots	تعداد نقاط مسطح
Cpoints	تعداد دفعاتی که سری زمانی از خط مبدا عبور می کند
Klscore	امتیاز Kullback-Leibler
Change.idx	محل داده با بیشترین امتیاز Kullback-Leibler

۳-۱-۱) خوشه بندی سری زمانی

هدف از این گام شناسایی خوشه های سری های زمانی همگن است. در این مرحله بردارهای ویژگی محاسبه شده برای سری های زمانی به عنوان ورودی الگوریتم خوشه بندی سری ها مورد استفاده قرار می گیرد. شناسایی خوشه های دارای سری های زمانی شبیه به هم، می تواند عملکرد مدل پیش بینی را بهبود بخشد.

الگوریتم ۱. خوشه‌بندی سری‌های زمانی براساس بردار ویژگی $\mathcal{F}$ FeaturesClustering ($\mathcal{F}, n$) $\mathcal{C} \leftarrow$ Cluster $\mathcal{F}$ using a clustering algorithm into n groupsReturn $\mathcal{C}$

۳-۱-۲ مرحله پیش‌بینی

پس از شناسایی خوشه‌ها، به ازای هر خوشه، یک GFM با استفاده از داده‌های سری زمانی آن خوشه آموزش داده می‌شود. مرحله پیش‌بینی شامل پیش‌پردازش سری‌های زمانی، استفاده از مدل پیش‌بینی برای آموزش GFM و برخی مراحل پس از پردازش است. علاوه بر الگوریتم‌های ۴-۲ که مراحل اصلی مرحله پیش‌بینی را نشان می‌دهد، جزئیات بیشتری را در زیر ارائه می‌دهیم.

الگوریتم ۲. مرحله پیش‌پردازش از پیش‌بینی مبتنی بر خوشه‌بندی

Forecasting ($\mathbf{D}, \mathcal{C}$, lag, look forward, frequency)

{Time series preprocessing stage}

Foreach time series $X_i \in \mathbf{D}$ do $m_i \leftarrow \text{Mean}(X_i)$ $X_i \leftarrow \frac{X_i}{m_i}$ # divide by mean $X_i^{\text{transformed}} \leftarrow \text{LogTransformation}(X_i)$ $\mathbf{D}^{\text{transformed}} \leftarrow \text{Append}(X_i^{\text{transformed}})$ # Add transformed data to list

End

Seasonalities $\leftarrow \text{DecomposeSeasonalities}(X_i^{\text{transformed}} \in \mathbf{D}^{\text{transformed}})$

الگوریتم ۳. مرحله ساخت مدل پیش‌بینی

{Model building stage}

HyperparametersGrid $\leftarrow$ Define Hyperparameter Grid for predictors $\theta_p: [\theta_p^1, \theta_p^2, \dots, \theta_p^n]$ Initializing the predictors corresponding to hyperparameters $\mathbf{M}$ Foreach cluster C_j in $\mathcal{C}$, j in $\{1, 2, \dots, n\}$ doForeach $X_i^{\text{transformed}}$ in C_j do

multi-input-multi-output (MIMO)

GeneratedSamples $\leftarrow$ GeneratingMIMO($X_i^{\text{transformed}}$, InWindow, OutWindow)

End

 $\mathbf{D}_{\text{train}}, \mathbf{D}_{\text{validation}}, \mathbf{D}_{\text{test}} \leftarrow \text{Splitting}(\text{GeneratedSamples})$ Best_GFM $_j \leftarrow \text{Train}(\mathbf{M}, \mathbf{D}_{\text{train}}, \mathbf{D}_{\text{validation}})$

End

الگوریتم ۴. مرحله پس‌پردازش

{Time series postprocessing stage}

Foreach cluster C_j in $\mathcal{C}$, j in $\{1, 2, \dots, n\}$ Foreach $X_i^{\text{transformed}}$ in C_j do $r_i \leftarrow \text{Predict}(\text{Best_GFM}_j, X_i^{\text{transformed}})$ $r_i^{\text{reseasonalize}} \leftarrow \text{Reseasonalization}(r_i)$ $r_i^{\text{denormalize}} \leftarrow \text{Denormalization}(r_i^{\text{reseasonalize}})$ $r_i^{\text{InverseLog}} \leftarrow \text{InverseLogTransform}(r_i^{\text{denormalize}})$ $\mathcal{R}_{C_j} \leftarrow \text{Append}(r_i^{\text{InverseLog}})$ #Add forecasting results to $\mathcal{R}^k$

End

 $\mathcal{R} \leftarrow \text{Append}(\mathcal{R}_{C_j})$

End

Return $\mathcal{R}$

به جهت اینکه سری‌های موجود در یک مجموعه داده سری زمانی، دارای مقیاس و روند مختلفی هستند، از این رو به منظور آموزش مناسب GFM روی $\text{TS} = \{X_i\}_{i=1}^N$ نیاز است تا پیش‌پردازش‌هایی روی داده‌ها صورت گیرد. پیش‌پردازش‌های صورت گرفته به این شرح است:

۳-۱-۳) نرمال سازی

برای یکنواخت کردن مقیاس در سری های زمانی در ابتدا مقادیر هر سری بر میانگین آن تقسیم می شوند (رابطه (۲)).

$$X_{i,normalized} = \frac{X_i}{\frac{1}{K} \sum_{t=1}^K X_{i,t}} \quad (2)$$

که در آن K نشان دهنده تعداد نقاط سری زمانی X_i و $i \in \{1, 2, \dots, N\}$ است.

۳-۱-۴) تبدیل لگاریتمی

به منظور تثبیت واریانس هر سری زمانی، تبدیل لگاریتمی روی سری نرمال شده با استفاده از رابطه (۳) صورت می گیرد.

$$T_{i,normalized \& \log_trans} = \begin{cases} \log(T_{i,normalized}), & \min(TS) > 0; \\ \log(T_{i,normalized} + 1), & \min(TS) = 0, \end{cases} \quad (3)$$

۳-۱-۵) غیر فصلی سازی

در این مرحله عمل غیر فصلی سازی روی سری زمانی انجام می شود. روش استخراج الگوی فصلی که از آن بهره خواهیم برد، معروف به تجزیه الگوی فصلی - روند با استفاده از الگوریتم LOESS موسوم به STL خواهد بود (کلوند^{۳۲} و همکاران، ۱۹۹۰). بطور خلاصه این روش به عنوان یک روش قوی برای تجزیه یک سری زمانی به سه بخش الگوی فصلی، الگوی روند و سری زمانی باقیمانده شناخته می شود. برای غیر فصلی سازی از رابطه (۴) برای تجزیه سری به مولفه های فصلی، روند، باقیمانده استفاده کرده و سپس با استفاده از رابطه (۵) مولفه فصلی را حذف می کنیم.

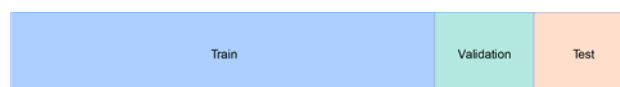
$$T_{i,normalized \& \log_trans} = S_i + t_i + r_i \quad (4)$$

$$T_{i,normalized \& \log_trans \& deseasonalized} = T_{i,normalized \& \log_trans} - S_i \quad (5)$$

در روابط (۴) و (۵) S_i ، t_i ، مولفه فصلی، t_i مولفه روند و r_i باقی مانده را نشان می دهد.

۳-۱-۶) تقسیم داده ها

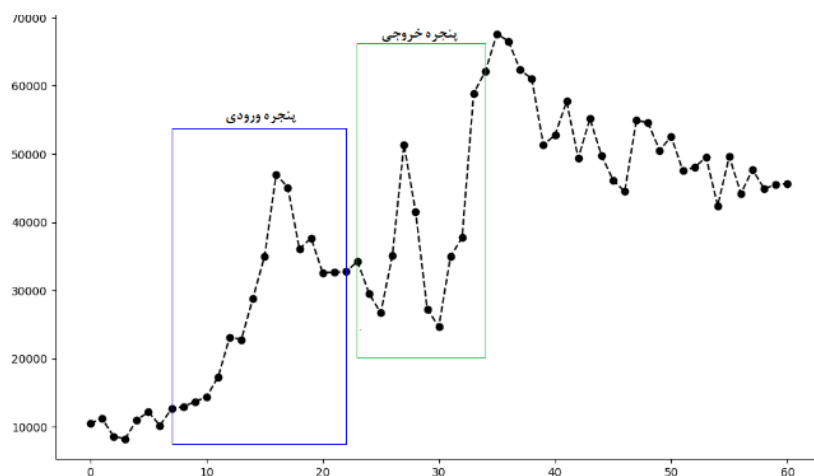
در مسئله پیش بینی سری زمانی، هدف پیش بینی h نقطه آینده یک سری است. بدین منظور در ابتدا داده های هر سری زمانی به سه بخش آموزشی، اعتبارسنجی و آزمون تقسیم می شود. بخش آموزشی و اعتبارسنجی در ساخت مدل پیش بینی مورد استفاده قرار می گیرند و بخش آزمون برای ارزیابی عملکرد مدل استفاده می شود (شکل ۲).



شکل ۲. مثالی از تقسیم داده های یک سری زمانی به سه بخش آموزشی، اعتبارسنجی و آزمون

۳-۱-۷) نمونه‌سازی با روش پنجره لغزان

بسیاری از مطالعات، مدل‌های پیش‌بینی را بر مبنای استراتژی پیش‌بینی تک‌نقطه طراحی می‌کنند. در این تحقیق برای پیش‌بینی h نقطه آینده، از استراتژی پیش‌بینی نقاط چندگانه MIMO استفاده می‌کنیم که در آن h نقطه آینده بطور همزمان توسط مدل پیش‌بینی می‌شود. همچنین برای تبدیل سری‌های زمانی به فرمت قابل استفاده توسط شبکه‌های عصبی از روشی که تحت عنوان پنجره لغزان است، استفاده می‌شود. در ابتدا یک سری زمانی با طول مشخص $tsLength$ با استفاده از این روش و به صورت یک گام در هر لغزش پنجره، به مجموعه‌ای از تکه‌هایی به تعداد $inputLength$ ، $tsLength$ و $outputLength$ تبدیل می‌شود. در اینجا $outputLength$ نشان‌دهنده افق پیش‌بینی خواهد بود و $inputLength$ نشانگر تعداد مقادیری است که به عنوان ورودی استفاده خواهند شد. در شکل ۳ نمونه‌ای از اعمال پنجره لغزان بر روی یک سری زمانی آورده شده است.



شکل ۳. نمونه‌ای از اعمال پنجره لغزان بر روی یک سری زمانی

۳-۱-۸) آموزش مدل پیش‌بینی

در این مرحله به ازای هر خوشه به دست آمده در مرحله پیش‌پردازش، یک مدل پیش‌بینی سراسری ساخته می‌شود. در مطالعات پیشین، گونه‌های مختلف شبکه عصبی بازگشتی مانند شبکه حافظه کوتاه مدت طولانی بطور گسترده مورد استفاده قرار گرفته است. برای بررسی اثربخشی روش پیشنهادی، ما مدل‌های پیش‌بینی مبتنی بر یادگیری عمیق را با استفاده از شبکه‌های کانولوشنی زمانی (TCN) پیشنهاد می‌کنیم (بای^{۳۳} و همکاران، ۲۰۱۸). TCN یک توسعه یافته از CNN است که در مقایسه با RNNها عملکرد بهتری نشان داده است. همچنین TCN قادر به پردازش موازی داده‌ها است که منجر به آموزش سریع تر مدل می‌شود. معماری شبکه TCN-CNN مورد استفاده در جدول ۲ توصیف شده است. همچنین برای ارزیابی بیشتر رویکرد پیشنهادی، یک روش مبتنی بر مکانیزم توجه چندگانه (وسوانی و همکاران^{۳۴}، ۲۰۱۷) را پیاده‌سازی می‌کنیم. این مدل شامل لایه‌های ورودی، مکانیزم توجه چندگانه، مسطح‌سازی^{۳۵}، کاملاً متصل^{۳۶} و خروجی است.

³³ Bai

³⁴ Vaswani

³⁵ Flatten

³⁶ Dense

جدول ۲. توپولوژی مدل TCN-CNN

توپولوژی مدل TCN-CNN
Input (shape: (1, Lag (input window)))
TCN (filters: 64, kernel size: 3, dilations: (1, 2, 4, 8))
1D Convolutional (filters: 64, kernel size: 4, activation: linear)
1D Convolutional (filters: 16, kernel size: 4, activation: linear)
Flatten
Dense (units: Forecast horizon, activation: linear)
Dense

۳-۲) پس پردازش

بعد از آموزش مدل پیش بینی، مدل به دست آمده روی بخش آموزش هر سری زمانی در مجموعه داده، اعمال می شود. مقادیر پیش بینی به دست آمده باید پس پردازش شوند تا به مقیاس اولیه داده ها برگردند. مراحل پس پردازش، معکوس مراحل پیش پردازش سری زمانی می باشد.

۳-۲-۱) فصلی سازی

به جهت اینکه در مرحله پیش پردازش، مقادیر مولفه های فصلی برای بخش آزمون هر سری زمانی (h نقطه پیش بینی) قابل محاسبه نبود، برای به دست آوردن مولفه های فصلی بخش آزمون، یک مدل هموارسازی نمایی را روی داده های مولفه فصلی بخش آموزشی هر سری زمانی آموزش می دهیم و از آن مدل برای پیش بینی مولفه های فصلی بخش آزمون استفاده می کنیم. در فصلی سازی، مولفه های فصلی به دست آمده به مقادیر پیش بینی هر سری زمانی افزوده می شوند.

۳-۲-۲) تبدیل معکوس لگاریتمی

در این مرحله روی مقادیر پیش بینی به دست آمده، تبدیل معکوس لگاریتمی انجام می شود.

۳-۲-۳) غیر نرمال سازی

در این مرحله، مقادیر پیش بینی به دست آمده در مقدار میانگین هر سری زمانی ضرب می شود و مقادیر پیش بینی نهایی به دست می آید.

۴) یافته ها

در این بخش ابتدا مجموعه داده مورد استفاده را توصیف می کنیم سپس به بیان تنظیمات آزمایشات، ابر پارامترها و تعریف معیارهای خطا پرداخته می شود. همچنین گونه های مختلف مدل های پیش بینی و مدل های محک را توصیف می کنیم. در نهایت عملکرد مدل های بکار گرفته شده را در مقایسه با مدل های معیار ارزیابی می کند.

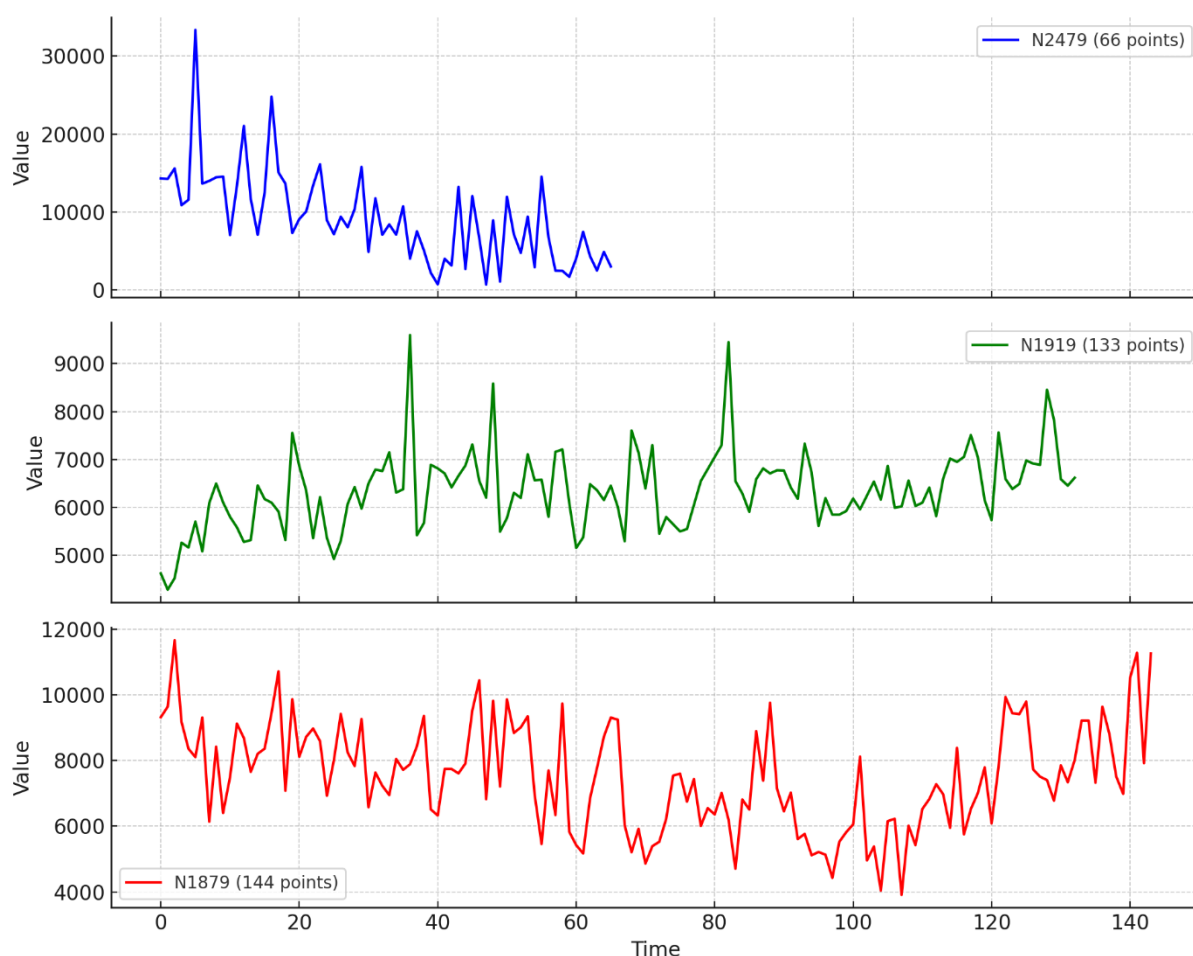
۴-۱) مجموعه داده

در این تحقیق، ما از مجموعه داده M3 برای ارزیابی روش پیشنهادی استفاده می کنیم. این مجموعه داده یکی از داده های مورد استفاده در مسابقات پیش بینی سری زمانی است. سه سری زمانی از این مجموعه داده با طول های متفاوت در شکل ۴ به تصویر کشیده شده است. M3 دارای ۶ مجموعه داده است که مشخصات آنها در جدول ۳ نمایش داده

شده است. این مجموعه داده دارای سری های زمانی با طول های غیریکسان است چنانکه طول کوتاهترین سری و بلندترین سری در این مجموعه به ترتیب برابر با ۶۶ و ۱۴۴ است. در تمام مدل های مورد استفاده، ما یک استراتژی پیش بینی چند گامی را اجرا کردیم که امکان پیش بینی برای h نقاط زمانی را بطور همزمان فراهم می کند. جایگاه h با افق پیش بینی در هر مجموعه داده مطابقت دارد. ویژگی های این مجموعه داده در جدول ۳ توصیف شده است.

جدول ۳. توصیف مجموعه داده M3

فرکانس	افق پیش بینی	تعداد سری	مجموعه داده
ماهانه	18	111	M3-demo
ماهانه	18	145	M3-finance
ماهانه	18	334	M3-industry
ماهانه	18	312	M3-macro
ماهانه	18	474	M3-micro
ماهانه	18	52	M3-other



شکل ۴. مصورسازی سه سری زمانی از مجموعه داده M3 با طول های متفاوت

۴-۲) تنظیمات آزمایشات و ابرپارامترها:

روش پیشنهادی با زبان پایتون و با استفاده از کتابخانه یادگیری عمیق Keras پیاده سازی شده است. همچنین برای استخراج ویژگی از سری های زمانی، از کتابخانه tsfeature در زبان R استفاده شده است. در مرحله پیش بینی، مهمترین ابرپارامترها اندازه پنجره ورودی، اندازه دسته، تعداد دور و نرخ یادگیری است. در جدول ۴ لیست محدوده مقادیر برای این ابرپارامترها نشان داده شده است. با توجه به غیر یکسان بودن طول سری های زمانی، اندازه پنجره ورودی برای هر مجموعه داده طوری انتخاب گردید که امکان استفاده حداکثری از اطلاعات گذشته سری زمانی وجود داشته باشد.

جدول ۴. محدوده مقادیر مورد استفاده ابرپارامترها

مجموعه داده	اندازه دسته	تعداد دورها	نرخ یادگیری
M3-Demo	[20-100]	[5-50]	0.01-0.0001
M3-Finance	[20-100]	[5-50]	0.01-0.0001
M3-Industry	[20-100]	[5-10]	0.01-0.0001
M3-Macro	[20-100]	[5-50]	0.01-0.0001
M3-Micro	[40-100]	[5-35]	0.01-0.0001
M3-Other	[20-100]	[5-50]	0.01-0.0001

۴-۳) معیارهای ارزیابی

در این مطالعه از دو معیار اصلی و معروف در ارزیابی مدل های پیش بینی سری های زمانی شامل sMAPE و MASE استفاده می کنیم (باندرا، ۲۰۲۳). با در نظر گرفتن اینکه در هر کدام از سری های زمانی h نقطه آخر هر سری به عنوان بخش آزمون در نظر گرفته می شود، این معیارها به صورت زیر تعریف می شوند.

$$sMAPE = \frac{200}{h} \sum_{i=1}^h \left(\frac{|A_i - F_i|}{|A_i| + |F_i|} \right) \quad (6)$$

$$MASE = \frac{1}{h} \frac{\sum_{i=1}^h (|A_i - F_i|)}{\frac{1}{n-M} \sum_{i=S+1}^n (|A_i - A_{i-S}|)} \quad (7)$$

در روابط (۶) و (۷) A_i و F_i نشان دهنده مقادیر واقعی و پیش بینی شده در نقطه i هستند. پارامتر h نشان دهنده افق پیش بینی است. n در رابطه (۷) نشان دهنده تعداد نقاط بخش آموزشی هر سری زمانی است. S دوره فصلی هر سری را نشان می دهد. برای یک مجموعه داده سری زمانی، هر کدام از این معیارهای خطا در سطح هر سری محاسبه شده و در نهایت میانگین خطاها به عنوان معیار نهایی کل مجموعه داده محاسبه می گردد و در نتایج گزارش ارائه می شود.

۴-۴) مدل های پیاده سازی شده:

با توجه به روش پیشنهادی و نوع الگوریتم خوشه بندی و همچنین نوع پیش پردازش سری ها، ۳ مدل پیش بینی مبتنی بر روش TCN-CNN و همچنین ۳ مدل مبتنی بر روش مکانیزم توجه چندگانه پیاده سازی گردید (جدول ۵). دو الگوریتم K-Medoids و خوشه بندی طیفی در مرحله خوشه بندی مورد استفاده قرار گرفتند. پارامتر تعداد خوشه ها برابر ۲ در نظر گرفته

شد. در تمامی مدل‌های پیش‌بینی از تابع زیان^{۳۷} MSE و همچنین بهینه‌ساز Adam^{۳۸} (کینگما^{۳۹} و همکاران) استفاده گردید. تعداد تکرارهای شبکه^{۴۰} با استفاده از روش توقف زودهنگام^{۴۱} مشخص می‌شود.

جدول ۵. مدل‌های پیاده‌سازی شده

نام مدل	الگوریتم خوشه‌بندی
TC-Baseline	عدم خوشه‌بندی
TC-CL-KM	K-Medoids
TC-CL-CP	خوشه‌بندی طیفی
MHA-Baseline	عدم خوشه‌بندی
MHA-CL-KM	K-Medoids
MHA-CL-CP	خوشه‌بندی طیفی

۴-۵) نتایج و تحلیل

نتایج مدل‌های پیاده‌سازی شده بر مبنای معیارهای میانگین sMAPE و MASE به ترتیب در جدول ۶ و ۷ آورده شده‌است. مطابق جدول ۷، تحلیل نتایج مدل‌های مبتنی بر TCN-CNN نشان می‌دهد که در معیار MASE مدل TC-CL-KM عملکرد بهتری را نسبت به مدل پایه و مدل TC-CL-CP از خود نشان داده‌است. همچنین بطور کلی میانگین خطای MASE برای مدل‌های مبتنی بر خوشه‌بندی کمتر از مدل پایه است که نشان‌دهنده اثربخشی بکارگیری روش‌های خوشه‌بندی در شناسایی سری‌های مرتبط به هم و ساخت مدل‌های خاص هر خوشه است. همچنین مطابق جدول ۷، تحلیل نتایج مدل‌های مبتنی بر مکانیزم توجه چندگانه بر مبنای معیار MASE حاکی از آن است که هر دو مدل مبتنی بر خوشه‌بندی، نتایج بهتری نسبت به مدل پایه متناظر خود (MHA-Baseline) کسب کرده‌اند. همچنین مقایسه کلی دو مدل TCN-CNN و مدل توجه چندگانه حاکی از برتری TCN-CNN در تمامی حالت‌ها است.

جدول همانطور که در جدول ۶ مشخص است، با بکارگیری مدل TCN-CNN به‌عنوان مدل پیش‌بینی، هر دو مدل مبتنی بر خوشه‌بندی در میانگین کلی نتایج بهتری نسبت به مدل پایه (TC-Baseline) در معیار خطای sMAPE دارند. مدل TC-CL-KM که از روش خوشه‌بندی K-Medoids بهره می‌گیرد، در معیار sMAPE در میانگین کلی بهتر از مدل TC-CL-CP عمل می‌کند. این مدل در دو مجموعه داده M3-macro و M3-Micro بهترین عملکرد را دارد و می‌توان گفت با توجه به اندازه این مجموعه داده، الگوریتم K-Medoids برای مجموعه داده‌های با اندازه بزرگتر بهتر از الگوریتم خوشه‌بندی طیفی عمل کرده‌است.

همچنین نتایج جدول ۶ نشان می‌دهد که با بکارگیری مدل توجه چندگانه (MHA) به‌عنوان مدل پیش‌بینی، هر دو مدل مبتنی بر خوشه‌بندی نسبت به مدل پایه متناظر خود (MHA-Baseline) در معیار sMAPE نتایج بهتری از خود نشان داده‌اند. مقایسه نتایج روش‌های مبتنی بر TCN-CNN و روش مکانیزم توجه چندگانه نشان می‌دهد که روش TCN-CNN

³⁷ Loss function

³⁸ Adaptive Moment Estimation

³⁹ Kingma

⁴⁰ Epoch

⁴¹ Early stopping

دارای عملکرد بهتری نسبت به روش مکانیزم توجه چندگانه در تمامی مدل ها دارد. به عبارت دیگر، هم در رویکرد بدون خوشه بندی و هم در رویکرد با خوشه بندی، روش TCN-CNN نتایج بهتری از خود نشان می دهد. دلیل این تفاوت می تواند ناشی از ماهیت رویکرد پیش بینی سراسری باشد که در آن نمونه های داده از مجموعه ای از سری های زمانی با خصوصیات متفاوت استخراج می شوند و این کار مانع یادگیری موثر روش مبتنی بر توجه می شود. شایان ذکر است که طبق تحقیقات قبلی، در پیش بینی سری های زمانی به صورت مجزا، مدل های پیش بینی مبتنی بر مکانیزم توجه، نتایج بهتری از خود نشان داده اند.

جدول ۶. نتایج مدل های پیاده سازی شده بر مبنای معیارهای SMAPE

مدل	M3-Demo	M3-Finance	M3-Industry	M3-Macro	M3-Micro	M3-Other	میانگین کل
نتایج مدل TCN-CNN							
TC-Baseline	8.62	14.43	13.03	6.88	21.20	10.04	14.09
TC-CL-KM	9.02	14.09	13.08	6.47	21.02	9.79	13.94
TC-CL-CP	8.58	13.76	13.13	6.67	21.21	9.73	13.99
نتایج مدل MHA							
MHA-Baseline	8.81	16.84	13.51	6.91	21.72	11.55	14.69
MHA-CL-KM	8.99	16.61	13.16	6.83	21.61	10.59	14.51
MHA-CL-SP	8.60	15.90	13.17	6.46	21.67	10.43	14.35

مطابق جدول ۷، تحلیل نتایج مدل های مبتنی بر TCN-CNN نشان می دهد که در معیار MASE مدل TC-CL-KM عملکرد بهتری را نسبت به مدل پایه و مدل TC-CL-CP از خود نشان داده است. همچنین بطور کلی میانگین خطای MASE برای مدل های مبتنی بر خوشه بندی کمتر از مدل پایه است که نشان دهنده اثربخشی بکارگیری روش های خوشه بندی در شناسایی سری های مرتبط به هم و ساخت مدل های خاص هر خوشه است. همچنین مطابق جدول ۷، تحلیل نتایج مدل های مبتنی بر مکانیزم توجه چندگانه بر مبنای معیار MASE حاکی از آن است که هر دو مدل مبتنی بر خوشه بندی، نتایج بهتری نسبت به مدل پایه متناظر خود (MHA-Baseline) کسب کرده اند. همچنین مقایسه کلی دو مدل TCN-CNN و مدل توجه چندگانه حاکی از برتری TCN-CNN در تمامی حالت ها است.

جدول ۷. نتایج مدل های پیاده سازی شده بر مبنای معیار MASE

مدل	M3-Demo	M3-Finance	M3-Industry	M3-Macro	M3-Micro	M3-Other	میانگین کل
نتایج مدل TCN-CNN							
TC-Baseline	0.916	1.238	1.034	1.003	0.715	0.619	0.918
TC-CL-KM	0.896	1.210	1.037	0.953	0.690	0.591	0.894
TC-CL-CP	0.833	1.162	1.051	0.988	0.698	0.594	0.898
نتایج مدل MHA							
MHA-Baseline	1.434	1.484	1.142	1.075	0.705	0.723	1.21
MHA-CL-KM	1.253	1.454	1.067	1.065	0.710	0.694	0.99
MHA-CL-CP	1.042	1.393	1.095	0.981	0.712	0.676	0.95

۴-۶) مقایسه با مدل‌های محک

در بخش قبلی به این نتیجه رسیدیم که مدل‌های مبتنی بر خوشه‌بندی نسبت به مدل‌های بدون خوشه‌بندی، عملکرد بهتری با هر دو مدل پیش‌بینی از خود نشان می‌دهند. در این بخش به منظور ارزیابی کامل روش پیشنهادی به مقایسه برترین مدل‌های به‌دست‌آمده در این تحقیق، یعنی مدل‌های TCN-CNN با مدل‌های محک می‌پردازیم. مدل‌های محک مورد استفاده شامل: مدل‌های گروهی FFORMA (مونتر و همکاران، ۲۰۲۰)، مدل CatBoost (پروخورنکوا^{۴۲} و همکاران ۲۰۱۸) هستند. همچنین مدل‌های مبتنی بر یادگیری عمیق شامل: مدل DeepAR (سالیانس و همکاران، ۲۰۲۰) و مدل N-BEATS (اورشکین^{۴۳} و همکاران، ۲۰۱۹) هستند. به جهت اینکه در مطالعات گذشته، یک خطای کلی برای مجموعه داده M3 ارائه شده است، برای مقایسه صحیح مدل‌های ارائه شده، در این بخش میانگین خطاها برای کل سری‌های موجود در مجموعه داده M3 که شامل ۱۴۲۶ مورد می‌باشد، در جدول ۸ و شکل‌های ۵ و ۶ گزارش می‌شود. در معیار sMAPE مدل‌های ارائه شده در این تحقیق نتایج بهتری نسبت به مدل‌های تحقیقات پیشین دارند. از منظر معیار MASE هر دو مدل مبتنی بر خوشه‌بندی شامل مدل TC-CL-KM و همچنین مدل TC-CL-CP نتایج بهتری نسبت به مدل‌های FFORMA، CatBoost، Deep AR و N-BEATS دارند.

جدول ۸. مقایسه بهترین مدل به‌دست‌آمده با مدل‌های موجود استخراج شده از (گوداهیا و همکاران، ۲۰۲۱)

مدل	معیار sMAPE	معیار MASE
TC-CL-KM	13.94	0.894
TC-CL-CP	13.99	0.898
FFORMA	14.51	0.914
CatBoost	16.41	1.065
DeepAR	15.74	1.167
N-BEATS	14.76	0.934

۴-۷) پیچیدگی زمانی روش پیشنهادی

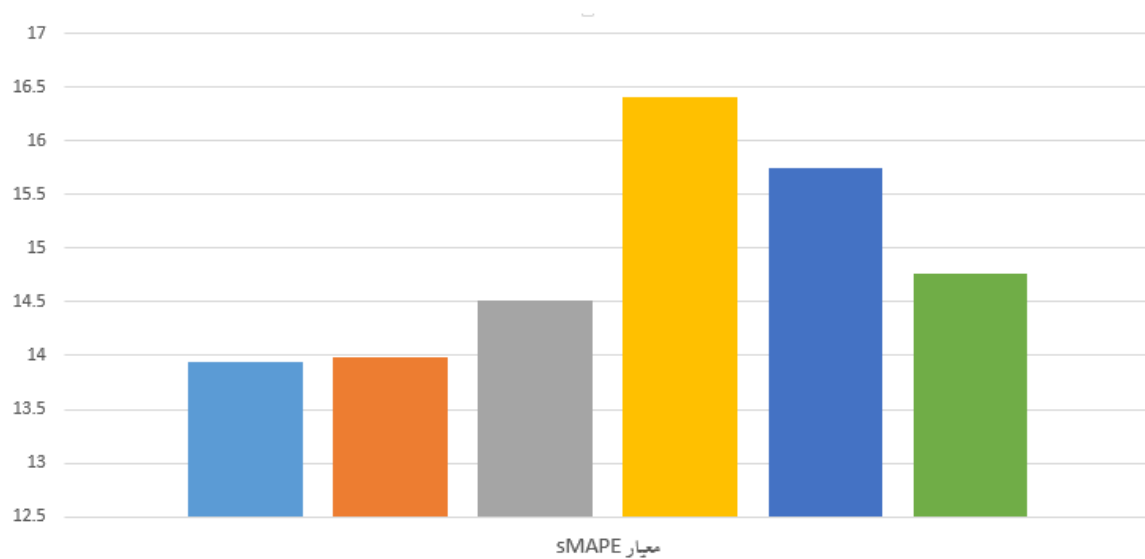
با توجه به اینکه در رویکرد پیشنهادی، سری‌های زمانی به خوشه‌های مختلف افراز شده و به ازای هر خوشه یک مدل سراسری خاص آموزش داده می‌شود، امکان موازی‌سازی آموزش مدل‌ها به وجود می‌آید. در جدول ۹ زمان اجرای هر مرحله از رویکرد پیشنهادی برای مدل پایه TC-Baseline و همچنین مدل مبتنی بر خوشه‌بندی TC-CL-CP بیان شده است. محیط محاسباتی مورد استفاده برای اجرای آزمایشات شامل منابع CPU و GPU است که شامل پردازنده Intel Xeon با ۲ هسته مجازی (vCPUs) و ۱۳ گیگابایت RAM و همچنین کارت گرافیک NVIDIA Tesla P100 با ۱۶ گیگابایت حافظه می‌باشد. همانطور که در این جدول مشخص است، زمان کل روش مبتنی بر خوشه‌بندی به دلیل موازی‌سازی کمتر از مدل پایه است که در آن یک مدل پیش‌بینی سراسری به ازای تمامی سری‌های زمانی در یک مجموعه آموزش داده می‌شود. بنابراین با در نظر گرفتن نتایج جداول ۶ و ۷ و همچنین تحلیل‌های محاسباتی، می‌توان به این نتیجه رسید که روش پیشنهادی موجب افزایش دقت و کاهش زمان اجرا می‌شود.

⁴² Prokhorenkova

⁴³ Oreshkin

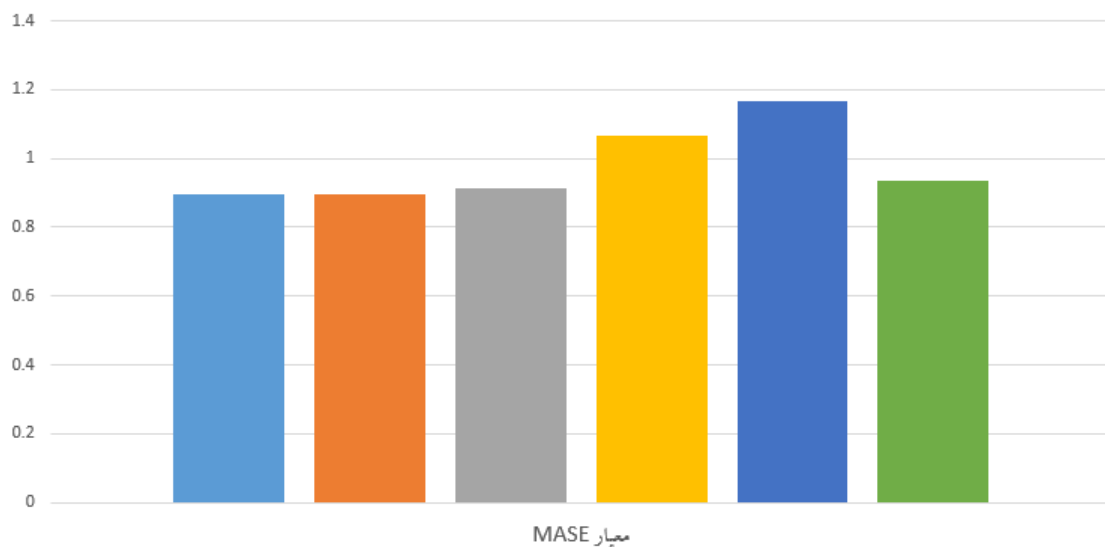
جدول ۹. زمان محاسباتی مدل مبتنی بر خوشه بندی و مدل پایه

مراحل رویکرد پیشنهادی	مدل	زمان (ثانیه)
استخراج ویژگی	TC-Baseline	-
	TC-CL-CP	27.15
خوشه بندی	TC-Baseline	-
	TC-CL-CP	0.2
پیش پردازش (مرحله پیش بینی)	TC-Baseline	59.9
	TC-CL-CP	25.44
آموزش مدل پیش بینی	TC-Baseline	325.78
	TC-CL-CP	122.76
زمان کل	TC-Baseline	385.68
	TC-CL-CP	175.55



■ TC-CL-KM ■ TC-CL-CP ■ FFORMA ■ CatBoost ■ DeepAR ■ N-BEATS

شکل ۵. مقایسه با مدل های محک از نظر معیار sMAPE



■ TC-CL-KM ■ TC-CL-CP ■ FFORMA ■ CatBoost ■ DeepAR ■ N-BEATS

شکل ۶. مقایسه با مدل‌های محک از نظر معیار MASE

(۵) نتیجه‌گیری و پیشنهادها

به دلیل تولید گسترده داده‌ها در قالب سری زمانی، مدل‌های پیش‌بینی سراسری نسبت به روش‌های پیش‌بینی سری‌های زمانی تک‌متغیره مورد توجه واقع شده‌اند. در این مطالعه، یک رویکرد مبتنی بر خوشه‌بندی، برای بهبود دقت مدل‌های سراسری به منظور مقابله با ناهمگونی داده‌ها در مجموعه‌های سری زمانی پیشنهاد داده شد. نتایج آزمایشات انجام شده بر روی مجموعه M3 نشان داد که در اکثر موارد، مدل‌های ایجاد شده با استفاده از روش پیشنهادی، عملکرد بهتری نسبت به مدل‌های پایه و همچنین مدل‌های محک از خود نشان داده‌اند. این امر قابلیت روش پیشنهادی را در شناسایی خوشه‌های همگن سری‌های زمانی و ساخت مدل خاص هر خوشه را نشان می‌دهد. علاوه بر این، تحلیل نتایج نشان داد که مدل‌های به دست آمده از رویکرد پیشنهادی و به ویژه مدل مبتنی بر خوشه‌بندی k-Medoids، میانگین sMAPE و میانگین MASE کمتری نسبت به مدل پایه که در آن خوشه‌بندی صورت نمی‌گیرد، دارند. همچنین مقایسه نتایج نشان داد مدل‌های حاصل از روش پیشنهادی، عملکرد بهتری نسبت به مدل‌های پیشرفته مانند DeepAR، N-BEATS، FFORMA و CatBoost دارند. مقایسه بیشتر با رویکردهای مبتنی بر خوشه‌بندی موجود در ادبیات موضوع نیز نشان داد مدل‌های پیشنهادی این تحقیق عملکرد بهتری از نظر میانگین sMAPE و میانگین MASE داشته‌اند.

آزمایشات ما نشان داد که خوشه‌بندی k-Medoids عملکرد بهتری نسبت به خوشه‌بندی طیفی داشت. عملکرد رویکرد پیشنهادی به استخراج ویژگی‌های معنی‌دار، خوشه‌بندی با کیفیت و همچنین مدل پیش‌بینی دقیق و مستحکم وابسته است. در تحقیقات آینده، قصد داریم روش پیشنهادی را با استفاده از شبکه‌های خودرمزنگار و روش‌های خوشه‌بندی عمیق توسعه دهیم.

منابع

- Bai, S., Kolter, J. Z., & Koltun, V. (2018). An empirical evaluation of generic convolutional and recurrent networks for sequence modeling. arXiv preprint arXiv:1803.01271. <https://doi.org/10.48550/arXiv.1803.01271>
- 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. <https://doi.org/10.1016/j.eswa.2019.112896>
- Bandara, K., Hewamalage, H., Liu, Y. H., Kang, Y., & Bergmeir, C. (2021). Improving the accuracy of global forecasting models using time series data augmentation. Pattern Recognition, 120, 108148. <https://doi.org/10.1016/j.patcog.2021.108148>
- Bandara, K. (2023). Forecasting with Big Data Using Global Forecasting Models. In Forecasting with Artificial Intelligence: Theory and Applications (pp. 107-122). M. Hamoudia, S. Makridakis, and E. Spiliotis, Editors. 2023, Cham: Springer Nature Switzerland: <https://doi.org/10.1007/978-3-031-35879-1>
- Box, G. E., Jenkins, G. M., Reinsel, G. C., & Ljung, G. M. (2015). Time series analysis: forecasting and control. John Wiley & Sons. <https://doi.org/10.1111/jtsa.12194>
- Cleveland, R. B., Cleveland, W. S., McRae, J. E., & Terpenning, I. (1990). STL: A seasonal-trend decomposition. J. off. Stat, 6(1), 3-73.
- Godahehwa, R., Bandara, K., Webb, G. I., Smyl, S., & Bergmeir, C. (2021). Ensembles of localised models for time series forecasting. Knowledge-Based Systems, 233, 107518. <https://doi.org/10.1016/j.knosys.2021.107518>
- Hewamalage, H., Bergmeir, C., & Bandara, K. (2021). Recurrent neural networks for time series forecasting: Current status and future directions. International Journal of Forecasting, 37(1), 388-427. <https://doi.org/10.1016/j.ijforecast.2020.06.008>
- Hewamalage, H., Bergmeir, C., & Bandara, K. (2022). Global models for time series forecasting: A simulation study. Pattern Recognition, 124, 108441. <https://doi.org/10.1016/j.patcog.2021.108441>

- Hyndman, R. J., Koehler, A. B., Snyder, R. D., & Grose, S. (2002). A state space framework for automatic forecasting using exponential smoothing methods. *International Journal of forecasting*, 18(3), 439-454. [https://doi.org/10.1016/S0169-2070\(01\)00110-8](https://doi.org/10.1016/S0169-2070(01)00110-8)
- Hyndman, R., Koehler, A. B., Ord, J. K., & Snyder, R. D. (2008). *Forecasting with exponential smoothing: the state space approach*. Springer Science & Business Media. <https://doi.org/10.1007/978-3-540-71918-2>
- Hyndman, R., Kang, Y., Montero-Manso, P., Talagala, T., Wang, E., Yang, Y., & O'Hara-Wild, M. (2019). tsfeatures: Time series feature extraction. R package version, 1(0).
- Januschowski, T., Gasthaus, J., Wang, Y., Salinas, D., Flunkert, V., Bohlke-Schneider, M., & Callot, L. (2020). Criteria for classifying forecasting methods. *International Journal of Forecasting*, 36(1), 167-177. <https://doi.org/10.1016/j.ijforecast.2019.05.008>
- Kingma, D. P., & Ba, J. (2014). Adam: A method for stochastic optimization. arXiv preprint arXiv:1412.6980.
- Laptev, N., Yosinski, J., Li, L. E., & Smyl, S. (2017, August). Time-series extreme event forecasting with neural networks at uber. In *International conference on machine learning* (Vol. 34, pp. 1-5). sn.
- Makridakis, S., Spiliotis, E., & Assimakopoulos, V. (2018). The M4 Competition: Results, findings, conclusion and way forward. *International Journal of forecasting*, 34(4), 802-808. <https://doi.org/10.1016/j.ijforecast.2018.06.001>
- Makridakis, S., Spiliotis, E., & Assimakopoulos, V. (2022). M5 accuracy competition: Results, findings, and conclusions. *International Journal of Forecasting*, 38(4), 1346-1364. <https://doi.org/10.1016/j.ijforecast.2021.11.013>
- Martínez, F., Frías, M. P., Pérez-Godoy, M. D., & Rivera, A. J. (2018). Dealing with seasonality by narrowing the training set in time series forecasting with kNN. *Expert systems with applications*, 103, 38-48. <https://doi.org/10.1016/j.eswa.2018.03.005>
- Montero-Manso, P., Athanasopoulos, G., Hyndman, R. J., & Talagala, T. S. (2020). FFORMA: Feature-based forecast model averaging. *International Journal of Forecasting*, 36(1), 86-92. <https://doi.org/10.1016/j.ijforecast.2019.02.011>
- Montero-Manso, P., & Hyndman, R. J. (2021). Principles and algorithms for forecasting groups of time series: Locality and globality. *International Journal of Forecasting*, 37(4), 1632-1653. <https://doi.org/10.1016/j.ijforecast.2021.03.004>
- Oreshkin, B. N., Carпов, D., Chapados, N., & Bengio, Y. (2019). N-BEATS: Neural basis expansion analysis for interpretable time series forecasting. arXiv preprint arXiv:1905.10437. <https://doi.org/10.48550/arXiv.1905.10437>
- Parmezan, A. R. S., Souza, V. M., & Batista, G. E. (2019). Evaluation of statistical and machine learning models for time series prediction: Identifying the state-of-the-art and the best conditions for the use of each model. *Information sciences*, 484, 302-337. <https://doi.org/10.1016/j.ins.2019.01.076>
- Prokhorenkova, L., Gusev, G., Vorobev, A., Dorogush, A. V., & Gulin, A. (2018). CatBoost: unbiased boosting with categorical features. *Advances in neural information processing systems*, 31. <https://doi.org/10.48550/arXiv.1706.09516>
- Salinas, D., Flunkert, V., Gasthaus, J., & Januschowski, T. (2020). DeepAR: Probabilistic forecasting with autoregressive recurrent networks. *International journal of forecasting*, 36(3), 1181-1191. <https://doi.org/10.1016/j.ijforecast.2019.07.001>
- Schubert, E., & Rousseeuw, P. J. (2021). Fast and eager k-medoids clustering: O(k) runtime improvement of the PAM, CLARA, and CLARANS algorithms. *Information Systems*, 101, 101804. <https://doi.org/10.1016/j.is.2021.101804>
- Smyl, S., & Kuber, K. (2016, June). Data preprocessing and augmentation for multiple short time series forecasting with recurrent neural networks. In *36th international symposium on forecasting*.
- Smyl, S. (2020). A hybrid method of exponential smoothing and recurrent neural networks for time series forecasting. *International journal of forecasting*, 36(1), 75-85. <https://doi.org/10.1016/j.ijforecast.2019.03.017>
- Vaswani, A., N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser and I. Polosukhin (2017). Attention is all you need. *Proceedings of the 31st International Conference on Neural Information Processing Systems*. Long Beach, California, USA, Curran Associates Inc.: 6000–6010.
- Ye, X., & Sakurai, T. (2016). Robust similarity measure for spectral clustering based on shared neighbors. *ETRI journal*, 38(3), 540-550. <https://doi.org/10.4218/etrij.16.0115.0517>