

UDC 004.89

DUAL-STREAM MODEL OF LONG SHORT-TERM MEMORY FOR AIR QUALITY PREDICTION USING STATION AND WEATHER DATA

Ivan Rudavskiy*  & Halyna Klym 

Lviv Polytechnic National University,
12 Stepan Bandera Str., Lviv, 79013, Ukraine

Rudavskiy, I., & Klym, H. (2026). Dual-Stream Model of Long Short-Term Memory for Air Quality Prediction Using Station and Weather Data. *Electronics and Information Technologies*, 34, 207–216. <https://doi.org/10.30970/eli.34.13>.

ABSTRACT

Background. Air pollution remains a critical environmental and public health issue, requiring accurate and reliable forecasting methods to support decision-making and mitigation strategies. With the increasing availability of environmental data from monitoring stations and meteorological sources, advanced data-driven approaches have become essential for modeling complex air quality dynamics. The development of robust deep learning architectures capable of capturing both temporal dependencies and external environmental influences is of significant importance.

Methods. This study proposes a Dual-Stream Long Short-Term Memory (LSTM) model for air quality forecasting, which processes station-based and meteorological data through two parallel streams. A preprocessing pipeline is applied, including K-nearest neighbors (KNN) imputation for handling missing values and Z-score normalization to standardize input features. The outputs of the LSTM streams are integrated using an attention-based fusion mechanism, which adaptively assigns importance to each data source. The model performance is compared with several baseline approaches.

Results and Discussion. Experimental results demonstrate that the proposed model achieves improved prediction accuracy and stability compared to traditional machine learning methods and single-factor models. The attention-based fusion mechanism effectively captures the dynamic relationships between station data and meteorological conditions, allowing the model to adapt to varying environmental influences. Evaluation using the Index of Agreement, coefficient of determination, and Root Mean Square Error confirms the superiority of the proposed approach in modeling complex air quality patterns.

Conclusion. The proposed Dual-Stream LSTM framework provides an effective solution for integrating heterogeneous environmental data in air quality forecasting tasks. By combining structured preprocessing techniques with an adaptive fusion strategy, the model enhances prediction performance and robustness. The results highlight the potential of multi-source deep learning approaches for environmental monitoring and can be extended to other time-series forecasting applications.

Keywords: air quality prediction; LSTM; dual-stream model; KNN imputation; Z-score; sensor.

INTRODUCTION

Air pollution remains one of the most critical environmental issues worldwide, posing significant threats to human health, ecosystems, and overall quality of life. As urbanization and industrial activities continue to intensify, the need for accurate and timely air quality forecasting becomes increasingly important for environmental monitoring, policy-making, and public health protection. Modern sensing infrastructures, supported by the rapid



© 2026 Ivan Rudavskiy & Halyna Klym. Published by the Ivan Franko National University of Lviv on behalf of Електроніка та інформаційні технології / Electronics and information technologies. This is an Open Access article distributed under the terms of the [Creative Commons Attribution 4.0 License](https://creativecommons.org/licenses/by/4.0/) which permits unrestricted reuse, distribution, and reproduction in any medium, provided the original work is properly cited.

development of Internet of Things (IoT) technologies, enable continuous collection of large-scale air quality and environmental data. However, despite the availability of such data, challenges related to data quality, heterogeneity, and complexity still limit the effectiveness of predictive models.

Air quality datasets are inherently influenced by multiple factors, including emission sources, meteorological conditions, and spatial interactions between monitoring stations. In practice, these datasets often suffer from noise, missing values, and inconsistencies caused by sensor malfunctions, communication failures, or environmental disturbances. Such imperfections can degrade model performance and reduce the reliability of forecasting results if not properly addressed during preprocessing and modeling stages.

A wide range of approaches has been explored for air quality prediction, including traditional statistical models [1, 2] and classical machine learning techniques [3, 4]. While these methods can provide reasonable results under certain conditions, they are generally limited in their ability to model complex nonlinear dependencies and long-term temporal patterns present in air pollutant concentrations, such as PM_{2.5} and PM₁₀. These limitations have motivated the adoption of deep learning methods, which are better suited for capturing intricate relationships in time series data.

Among deep learning approaches, Long Short-Term Memory networks have demonstrated strong performance in sequential prediction tasks due to their capability to retain long-term dependencies and model nonlinear temporal dynamics. This makes LSTM particularly effective for air quality forecasting, where pollutant levels are affected by both historical trends and external influences such as weather conditions. Nevertheless, many existing LSTM-based studies rely on single-source input data or combine multiple factors within a single model without explicitly distinguishing their individual contributions.

To better exploit the multi-source nature of air quality data, it is beneficial to consider separate modeling of different factor groups. In particular, station-based measurements primarily capture temporal dynamics of pollutant concentrations, while meteorological variables (e.g., temperature, humidity, wind speed) represent exogenous factors that significantly influence pollutant dispersion and formation. Treating these inputs independently allows for more specialized feature extraction and reduces the risk of information interference within a single model.

This paper proposes a Dual-Stream LSTM model for air quality forecasting, which processes station data and weather data through two parallel LSTM streams. Each stream is designed to learn distinct patterns associated with its respective input type, and their outputs are subsequently combined to generate the final prediction. This architecture enables more effective modeling of both temporal and exogenous dependencies while preserving the unique characteristics of each data source.

In this work, missing values are handled using a data-driven imputation strategy based on the K-nearest neighbors method, which estimates missing observations by exploiting similarities between data points without relying on strict distributional assumptions. Compared to classical statistical approaches, KNN imputation is more flexible and better suited for complex environmental datasets [5].

In addition, data normalization is applied to ensure stable and efficient model training. Specifically, Z-score normalization is used to standardize input features by removing the mean and scaling to unit variance. This step is particularly important in the proposed framework, as it allows heterogeneous data sources, such as pollutant concentrations and meteorological variables, to be processed in a consistent manner across the dual-stream architecture [6].

The performance of the proposed Dual-Stream LSTM model is evaluated against baseline methods, including traditional machine learning models and standard LSTM architectures. Experimental results demonstrate that the proposed approach more effectively captures complex relationships between air quality and influencing factors, leading to improved forecasting accuracy and robustness.

METHODS

In the course of the process of collecting air quality data, it is possible that monitoring stations may, on occasion, produce incomplete records. Such incompleteness may be due to malfunctions of the sensors, maintenance activities, or communication interruptions. If this issue is not addressed in a satisfactory manner, the accuracy and reliability of predictive models may be adversely affected [7]. To mitigate this issue, the K-nearest neighbors imputation method is used for data imputation.

The K-nearest neighbors imputation method is a non-parametric data reconstruction technique widely used to handle missing values in datasets with complex and unknown distributions. Unlike statistical approaches that rely on predefined assumptions, KNN imputes missing values by identifying the most similar observations in the dataset and using their information to estimate the missing entries. This data-driven strategy allows the method to preserve local structures and relationships within the data, making it particularly suitable for environmental and time-series datasets with nonlinear characteristics. [8]

The core idea of KNN imputation is based on the similarity measurement between data samples. For each instance containing missing values, a set of the k most similar instances (neighbors) is selected according to a distance metric, such as Euclidean distance, computed over the available features. The missing value is then estimated using an aggregation of the corresponding values from these neighbors, typically by calculating their mean or a distance-weighted average. As a result, the imputed values reflect patterns observed in similar data points, ensuring consistency with the underlying data structure.

The KNN imputation process consists of two main stages. In the first stage, distances between samples are computed using only the observed features, and the nearest neighbors are identified for each incomplete instance. In the second stage, the missing values are reconstructed by aggregating the values of the selected neighbors [9]. This approach enables flexible handling of different data types and does not require explicit modeling of the data distribution.

The mathematical formulation of the KNN imputation procedure is provided in [Eq. \(1\)](#), where the missing value is estimated as a function of the values of its nearest neighbors. This formulation allows the method to adapt to local data patterns and improves the overall completeness and reliability of the dataset for subsequent modeling tasks.

$$x_i^m = \frac{\sum_{j \in N_k(i)} (w_{ij} \cdot x_j^m)}{\sum_{j \in N_k(i)} w_{ij}}, \quad (1)$$

where $w_{ij} = [d(i, j) + \varepsilon]^{-1}$, $d(i, j)$ is the distance between sample i and neighbor j , ε is a small constant to avoid division by zero, x_j^m is the value of feature m in neighbor j .

Normalization Process: Air quality and meteorological variables often exhibit different scales and units, which can adversely affect the convergence and stability of machine learning models. To mitigate this issue, Z-score normalization is applied to standardize each feature by removing the mean and scaling to unit variance. [10] This transformation ensures that all input variables contribute proportionally during model training, preventing features with larger numerical ranges from dominating the learning process. Z-score normalization has been shown to improve model convergence, enhance stability, and increase the accuracy of forecasting predictions. The formula for Z-score normalization is presented in [Eq. \(2\)](#):

$$x'_i = x_i - \mu / \sigma, \quad (2)$$

where x_i is the original value of the feature, μ is the mean of the feature over the training data, σ is the standard deviation of the feature, x'_i is the normalized value.

To effectively integrate the information extracted from the two data streams, an attention-based fusion mechanism is employed. Instead of directly combining the outputs of the LSTM models, the proposed approach assigns adaptive importance weights to each stream, allowing the model to dynamically determine the contribution of station-based and meteorological features to the final prediction.

h_s and h_w denote the hidden representations produced by the LSTM models for station data and weather data, respectively. The attention mechanism computes a set of weights that reflect the relative importance of each representation. These weights are then used to construct a fused feature vector, which captures both temporal dependencies and external influences. The attention weights are calculated as presented in **Eq (3)**:

$$\begin{aligned} a_s &= \exp(e_s) / (\exp(e_s) + \exp(e_w)), \\ a_w &= \exp(e_w) / (\exp(e_s) + \exp(e_w)), \end{aligned} \quad (3)$$

where e_s and e_w are the attention scores for the station and weather streams, respectively. These scores are typically obtained through a learnable function, such as a fully connected layer:

$$e_s = w^T h_s, \quad e_w = w^T h_w, \quad (4)$$

where h_s and h_w are the hidden feature representations obtained from the LSTM models for station data and meteorological data, respectively, w^T is a learnable weight vector of the attention layer, e_s and e_w are the corresponding attention scores reflecting the importance of each data stream.

The fused feature representation h_f is then computed as a weighted combination of the two feature vectors:

$$h_f = \alpha_s h_s + \alpha_w h_w, \quad (5)$$

where h_s and h_w are the feature vectors from the station and meteorological streams, respectively, α_s and α_w are the normalized attention weights assigned to each stream.

The fused feature vector is passed through a fully connected layer to generate the final prediction output.

This attention-based fusion strategy enables the model to adaptively emphasize either temporal patterns from station data or external meteorological influences depending on the current environmental conditions. Compared to simple concatenation or static weighting methods, the proposed approach provides greater flexibility and improves the model's ability to capture complex relationships in air quality data.

RESULTS AND DISCUSSION

The workflow of the proposed LSTM-based air quality prediction is presented in **Fig. 1**.

Figure 2 shows a comparison of the predicted and measured PM_{2.5} concentrations, and **Figure 3** shows the corresponding comparison for PM₁₀ concentrations. The results demonstrate the effectiveness of the proposed method in modelling air pollutant behavior, as it produces predictions that closely match the observed values.

The Index of Agreement (IA) is a normalized statistical metric used to quantify the degree of correspondence between predicted and observed values. Its values range from 0 to 1, where 1 indicates perfect agreement between model outputs and observations, while values closer to 0 reflect poor predictive performance [11].

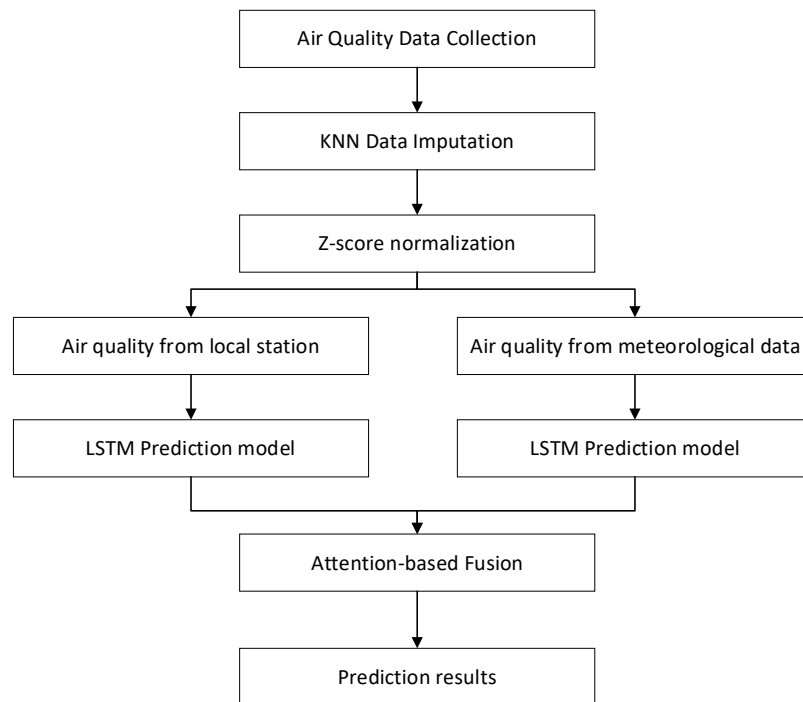


Fig 1. Workflow of the proposed LSTM-based air quality prediction

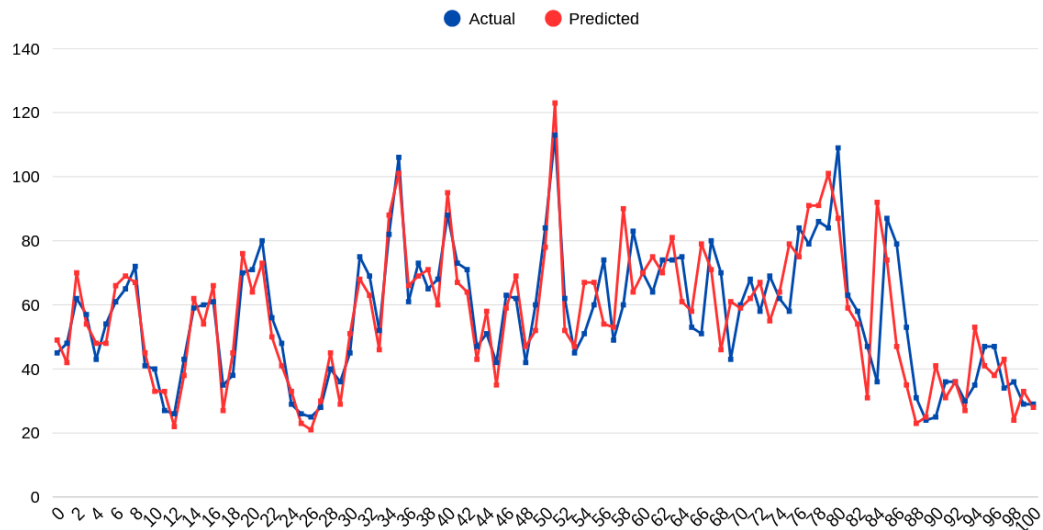


Fig 2. The comparison of predicted and actual values for the PM2.5 parameter

The IA evaluates model accuracy by comparing the magnitude of prediction errors to the total potential deviation in the observed data. Unlike simple error metrics, it accounts for both systematic and random discrepancies, providing a more comprehensive measure of model performance. This makes it particularly suitable for assessing environmental prediction models, where variability and uncertainty are inherent in the data.

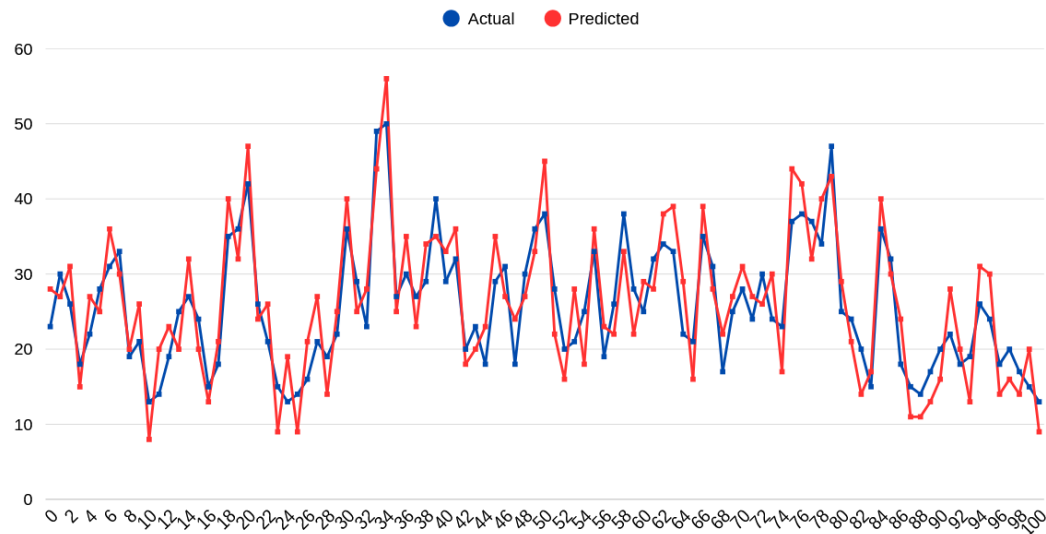


Fig 3. The comparison of predicted and actual values for the PM10 parameter

By incorporating both the variance of the observed values and the deviations of predictions from the mean, the Index of Agreement offers a balanced evaluation of how well the model reproduces observed patterns. As a result, it serves as a reliable indicator of the consistency and overall quality of forecasting results.

The mathematical formulation of the Index of Agreement is given in [Eq. \(6\)](#). [12]

$$IA = \sum_{i=1}^n (O_i - P_i)^2 / \sum_{i=1}^n (|P_i - \bar{O}| + |O_i - \bar{O}|)^2, \quad (6)$$

where, O_i is the observation value, P_i is the predicted value, $\bar{O}$ is the average observation value, $\bar{P}$ is the average predicted value.

Root mean square error (RMSE) is a widely adopted metric for evaluating the accuracy of predictive models, measuring the deviation between predicted and observed values. It is calculated as the square root of the average of the squared differences between the model's predictions and the actual observations. This provides an aggregate measure of the magnitude of prediction errors.

By squaring the individual errors before averaging them, the RMSE places greater importance on larger deviations, making it particularly sensitive to significant prediction errors. This characteristic is useful in applications such as air quality forecasting, where large discrepancies may have critical implications.

RMSE offers a straightforward measure expressed in the same units as the target variable, facilitating direct comparison between predicted and observed values. Consequently, it serves as an effective indicator of the model's overall predictive performance. The formula for RMSE calculation is defined in [Eq. 7](#).

$$RMSE = \sqrt{\sum_{i=1}^n (y_i - \hat{y}_i)^2 / n}, \quad (7)$$

where, y_i is the observation value, $\hat{y}_i$ is the predicted value, n is the number of observed values.

The coefficient of determination (CoD) R^2 is a widely used statistical measure that evaluates how well a regression model explains the variability in observed data. It represents the proportion of variance in the dependent variable that can be explained by the model. Values of R^2 range between 0 and 1; values closer to 1 indicate that the model successfully captures most of the underlying variability in the data, while values near 0 suggest limited explanatory power.

This metric provides insight into how well the predicted values align with actual observations by comparing residual errors with total data variation. As such, it is commonly used to evaluate the effectiveness and reliability of predictive models, particularly in regression-based forecasting tasks.

The coefficient of determination is useful for evaluating how well a model reproduces overall trends in the data, although it does not directly reflect the magnitude of prediction errors. Therefore, it is often used alongside other performance metrics to provide a more comprehensive evaluation. The formula for the CoD is defined in [Eq. 8](#).

$$R^2 = 1 - \left(\sum_{i=1}^n (y_i - \hat{y}_i)^2 / \sum_{i=1}^n (y_i - \bar{y}_i)^2 \right), \quad (8)$$

where, y_i is the observation value, $\hat{y}_i$ is the predicted value, $\bar{y}_i$ is the average value of y_i , n is the number of observed values.

Table 1 presents a comparative evaluation of three performance metrics – the Index of Agreement (IA), the CoD, and RMSE for the proposed LSTM-based model and several baseline approaches, including Support Vector Regression, XGBoost, Ridge Regression, and a Single-Factor Prediction model.

Table 1. Comparison of the prediction results of the different models

Model	IA	RMSE	CoD
Our model	0.92	11.7	0.77
Single-factor LSTM	0.88	26.14	0.63
XGBoost	0.37	67.94	0.75
Support Vector Regression	0.25	56.13	0.12
Ridge Regression	0.24	79.31	0.49

The results presented in [Table 1](#) demonstrate that the proposed Dual-Stream LSTM model provides the most balanced and accurate forecasting performance among the evaluated approaches. The improvement can be attributed to the ability of the model to simultaneously capture temporal dependencies in air quality measurements and the influence of meteorological conditions through the dual-stream architecture. Ridge Regression and Support Vector Regression exhibit substantially lower performance, which may be explained by their limited capacity to represent complex nonlinear relationships and long-term temporal patterns inherent in environmental data. XGBoost achieves a coefficient of determination comparable to that of the proposed model; however, the considerably lower Index of Agreement and higher RMSE values indicate reduced prediction consistency and larger forecasting errors. The comparison with the Single-Factor LSTM model demonstrates the contribution of meteorological information to the prediction process, as the integration of weather-related variables enables more comprehensive modeling of pollutant dynamics. The obtained results confirm the effectiveness of combining multi-source data with an attention-based fusion mechanism for air quality forecasting.

The model is trained using the Adam optimizer with a learning rate of 0.001. The sequence length (time step) is set to 24 to capture daily temporal patterns in air quality data. A batch size of 32 is used to ensure stable gradient updates.

The LSTM networks consist of 64 hidden units with a dropout rate of 0.2 to mitigate overfitting. The models are trained for 100 iterations. For the KNN imputation method, the number of nearest neighbors is set to $k = 5$, which provides a balance between local sensitivity and robustness.

Table 2 provides the training hyperparameters of the model.

Table 2. Training the hyperparameters of the model

Parameter	Value
Number of iterations	150
Time step	24
Batch size	32
Learning rate	0.001
Hidden units	64
k (for KNN imputation)	5

CONCLUSION

This paper presents a Dual-Stream LSTM framework for air quality forecasting that integrates information from monitoring stations and meteorological data sources. The main scientific contribution of the study lies in the combination of dual-stream temporal feature extraction with an attention-based fusion mechanism, enabling the model to dynamically determine the relative importance of different data sources during prediction. This approach allows more effective modeling of the complex interactions between air pollutant concentrations and environmental conditions than conventional single-source forecasting methods.

The experimental evaluation demonstrates that the proposed model achieves superior predictive performance compared with traditional machine learning approaches and single-factor prediction models. The results indicate that incorporating complementary information from meteorological variables significantly improves forecasting accuracy and model consistency, highlighting the benefits of multi-source learning for environmental prediction tasks.

The proposed framework can support intelligent air quality monitoring systems by providing more reliable forecasts for environmental management, public health protection, and urban planning. Furthermore, the developed methodology is not limited to air quality applications and may be adapted to other multivariate time-series forecasting problems involving heterogeneous data sources.

The proposed framework provides a flexible and efficient solution for modeling complex relationships in air quality data. By combining multi-source inputs with an adaptive fusion strategy, it offers improved predictive performance and can be extended to other environmental forecasting applications.

ACKNOWLEDGMENTS AND FUNDING SOURCES

This work was supported by the Ministry of Education and Science of Ukraine (Project No. 0125U001883).

COMPLIANCE WITH ETHICAL STANDARDS

The authors declare that they have no competing interests.

AUTHOR CONTRIBUTIONS

Conceptualization, [I.R., H.K.]; methodology, [I.R., H.K.]; investigation, [I.R., H.K.]; writing – original draft preparation, [I.R., H.K.]; writing – review and editing, [I.R., H.K.]; visualization, [I.R., H.K.].

All authors have read and agreed to the published version of the manuscript.

REFERENCES

- [1] Kotlia, P., Pant, J., & Lohani, M. C. (2025, September). Design and Implementation of Statistical Time Series Forecasting Model for Air Quality Assessment. In *2025 6th International Conference on Electronics and Sustainable Communication Systems (ICESC)* (pp. 913-918). IEEE. <https://doi.org/10.1109/icesc65114.2025.11212556>
- [2] Pandian, E., et al. (2025). Air quality forecasting for predictive analysis using machine learning. In *Proceedings of the 9th International Conference on Inventive Systems and Control (ICISC)* (pp. 187–192). IEEE. <https://doi.org/10.1109/icisc65841.2025.11188992>
- [3] Potey, S., et al. (2025). Air quality prediction using machine learning: A comprehensive review. In *Proceedings of the International Conference on Cognitive Computing in Engineering, Communications, Sciences and Biomedical Health Informatics (IC3ECBHI)* (pp. 1089–1093). IEEE. <https://doi.org/10.1109/ic3ecsbhi63591.2025.10991003>
- [4] Borah, J., et al. (2024). Timezone-aware auto-regressive long short-term memory model for multipollutant prediction. *IEEE Transactions on Systems, Man, and Cybernetics: Systems*, 1–9. <https://doi.org/10.1109/tsmc.2024.3463960>
- [5] Alani, N. H. S., Chand, P., & Al-Rawi, M. (2025). A Two-Stage Machine Learning Framework for Air Quality Prediction in Hamilton, New Zealand. *Environments*, 12(9), 336. <https://doi.org/10.3390/environments12090336>
- [6] Çelikten, H. (2026). A Multi-Output Deep Learning Framework for Simultaneous Forecasting of PM10 and Air Quality Index in High-Altitude Basins: A Case Study of Igdir, Türkiye. *Sustainability*, 18(8), 3883. <https://doi.org/10.3390/su18083883>
- [7] Rudavskiy, I., & Klym, H. (2026). Air quality prediction based on long short-term memory prediction model. *Measuring Equipment and Metrology*, 87(1). <https://doi.org/10.23939/istcmtm2026.01.026>
- [8] Zhang, X., Sun, Z., Zhou, Z., Jamali, S., & Liu, Y. (2022). Analysis and Dynamic Monitoring of Indoor Air Quality Based on Laser-Induced Breakdown Spectroscopy and Machine Learning. *Chemosensors*, 10(7), 259. <https://doi.org/10.3390/chemosensors10070259>
- [9] Taştan, M. (2025). Machine Learning–Based Calibration and Performance Evaluation of Low-Cost Internet of Things Air Quality Sensors. *Sensors*, 25(10), 3183. <https://doi.org/10.3390/s25103183>
- [9] Okorn, K., & Hannigan, M. (2021). Improving Air Pollutant Metal Oxide Sensor Quantification Practices through: An Exploration of Sensor Signal Normalization, Multi-Sensor and Universal Calibration Model Generation, and Physical Factors Such as Co-Location Duration and Sensor Age. *Atmosphere*, 12(5), 645. <https://doi.org/10.3390/atmos12050645>
- [11] Zhou, H., Wang, T., Zhao, H., & Wang, Z. (2022). Updated prediction of air quality based on kalman-attention-LSTM network. *Sustainability*, 15(1), 356. <https://doi.org/10.3390/su15010356>

- [12] Ma, X., Chen, T., Ge, R., Xv, F., Cui, C., & Li, J. (2023). Prediction of PM_{2.5} Concentration Using Spatiotemporal Data with Machine Learning Models. *Atmosphere*, 14(10), 1517. <https://doi.org/10.3390/atmos14101517>

ДВОПОТОКОВА МОДЕЛЬ ДОВГОЇ КОРОТКОЧАСНОЇ ПАМ'ЯТІ ДЛЯ ПРОГНОЗУВАННЯ ЯКОСТІ ПОВІТРЯ З ВИКОРИСТАННЯМ СТАНЦІЙНИХ ТА ПОГОДНИХ ДАНИХ

Іван Рудавський, Галина Клим

Національний університет «Львівська політехніка»,
вул. Степана Бандери 12, 79013 м. Львів, Україна

АНОТАЦІЯ

Вступ. Забруднення повітря залишається серйозною проблемою для навколишнього середовища та здоров'я населення, що вимагає точних і надійних методів прогнозування для підтримки процесу прийняття рішень та розробки стратегій щодо зменшення негативного впливу. Зі зростанням доступності екологічних даних, що надходять зі станцій моніторингу та метеорологічних джерел, сучасні підходи на основі даних стали незамінними для моделювання складної динаміки якості повітря. Велике значення має розробка надійних архітектур глибокого навчання, здатних враховувати як часові залежності, так і зовнішні екологічні впливи.

Методи. У цьому дослідженні пропонується модель довгої короткочасної пам'яті (ДКЧП) із двома потоками для прогнозування якості повітря, яка обробляє дані станцій та метеорологічні дані за допомогою двох паралельних потоків. Застосовується конвеєр попередньої обробки, що включає підстановку відсутніх значень методом k -найближчих сусідів (k -НС) та нормалізацію вхідних ознак за стандартизованою z -оцінкою. Вихідні дані потоків ДКЧП інтегруються за допомогою механізму злиття на основі уваги, який адаптивно присвоює важливість кожному джерелу даних. Ефективність моделі порівнюється з декількома базовими підходами.

Результати. Експериментальні результати свідчать про те, що запропонована модель забезпечує вищу точність прогнозування та стабільність порівняно з традиційними методами машинного навчання та однофакторними моделями. Механізм об'єднання даних на основі уваги ефективно враховує динамічні взаємозв'язки між даними станцій та метеорологічними умовами, що дозволяє моделі адаптуватися до мінливих впливів навколишнього середовища. Оцінка за допомогою індексу узгодженості, коефіцієнта детермінації та середньоквадратичної похибки.

Висновки. Запропонована архітектура двопотокової ДКЧП забезпечує ефективне рішення для інтеграції неоднорідних екологічних даних у завданнях прогнозування якості повітря. Завдяки поєднанню методів структурованої попередньої обробки даних із стратегією адаптивного об'єднання даних модель підвищує ефективність та надійність прогнозування. Отримані результати підкреслюють потенціал підходів на основі багатоджерельного глибокого навчання для моніторингу навколишнього середовища та можуть бути застосовані в інших задачах прогнозування часових рядів.

Ключові слова: Прогнозування якості повітря; ДКЧП; модель з двома потоками; метод k -НС; стандартизована z -оцінка; сенсор.