

Research Article

Maximum Normalized Orthogonal Projective Extreme Learning Machine for Green Energy Resource Data Analytics on Sustainable Environmental Air Pollution

R. Sudha Abirami¹, M.S. Irfan Ahmed²

^{1,2}Department of Computer Science & Research Centre, Thassim Beevi Abdul Kader College for Women, Kilakarai, Alagappa University, Karaikudi, Tamilnadu, India.

¹Corresponding Author : sudha.tbakc@gmail.com

Received: 25 August 2025

Revised: 21 May 2026

Accepted: 03 June 2026

Published: 28 July 2026

Abstract - Air pollution is a global environmental issues with severe threats and damage to human health and ecosystems. Air pollutants like Particulate Matter ($PM_{2.5}$), Nitrogen Dioxide (NO_2), Sulfur Dioxide (SO_2), Carbon Monoxide (CO), and Ozone (O_3) has been linked to adverse effects on the environment, including climate change and greenhouse gas emissions. Many researchers carried out their research on machine learning and artificial intelligence techniques for predicting environmental air pollution and for increasing the air quality. But the researcher's faces the problem of meeting the growing energy demand with minimum greenhouse gas emission and environmental vulnerability. Therefore, a large amount of air quality data is used for improving green energy efficiency. In this paper, the Maximum Normalized Orthogonal Lemma Projective Extreme Learning Machine Classification (MNOLPELMC) Method is introduced for sustainable environmental air pollution prediction with high accuracy and less time complexity. An experimental evaluation of the MNOLPELMC method is carried out with performance metrics such as prediction accuracy, precision, recall, root mean square error, and prediction time.

Keywords - Green energy resource data analytics, Sustainable environmental pollution, Extreme Machine Learning Model, Grubbs maximum normalized residual test, Random orthogonal lemma projection process.

1. Introduction

Green energy resource data analytics refers to the process of gathering, analyzing, and interpreting the data from different sources to optimize their performance, recognize potential issues, and make decisions for more efficient and sustainable energy operations. Many machine learning and deep learning models have been developed for forecasting of air pollutant concentration levels for effective air quality management and timely implementation.

A Gaussian-mixture Nested Factorial Variational Autoencoder (NF-VAE) was developed in [1] for multivariate air pollution prediction, achieving better accuracy and minimal root mean square error. However, the designed model failed to perform data preprocessing for enhancing its efficiency and minimizing the time.

A Neural Transfer learning model was designed in [2] with the aim of increasing the accuracy of air pollution prediction and reduced the mean square error. But it failed to focus on expanding the scope of analysis to include additional pollutants, as well as exploring the scalability of the approach in the complex environmental challenges. An artificial

intelligence and machine learning model was developed in [3]. But the integration of AI and ML with Numerical Weather Prediction (NWP) models was not developed to proactively manage air quality. An efficient machine learning model called Random Forest was developed in [4]. However, the Random Forest model caused more time complexity in air pollution prediction. Different machine learning models were developed in [5].

The designed model reduces the prediction time but fails to extend it by employing deep learning techniques for more accurate air quality index prediction. Machine learning algorithms were designed in [6]. However, it failed to use a novel method to successfully demonstrate the ad and further refine the approach to enhance the accuracy. The Machine Learning Linear Regression model was designed in [7].

However, the designed regression method failed to integrate accurate monitoring with a large volume of data samples. A novel Bayesian Optimization for Hybrid Time-Series Model with SARIMA called (BO-HyTS) was developed in [8]. However, the designed model failed to consider the multivariate meteorological data acquired from



multiple databases and design a more robust model for accurate prediction. A Data-Driven Supervised Machine Learning model was designed in [9].

But it dialed to involve a preprocessing stage for the input data to enhance the model performance. The Artificial Neural Networks (ANNs) model was designed in [10]. But the time complexity in the prediction was high. A multilayer perceptron neural network was introduced in [11]. However, it failed to perform significant feature selection to minimize the time consumption during prediction.

A Spatio-Temporal Air Quality Forecaster (ST-AQF) was developed in [12]. However, it did not utilize reliable, adaptive, robust, and flexible air quality forecasting models. An AI and advanced technology was developed in [13]. However, the performance analysis of accuracy and error remained unaddressed. To increase prediction accuracy, a new deep learning model was developed in [14].

However, it failed to optimize the network parameter settings for extracting spatio-temporal features for enhancing prediction performance. Machine learning methods were developed in [15]. But it failed to perform the analysis of air pollution in smart cities.

1.1. Major contributions

- To improve the prediction accuracy of air pollution, the MNOLPELMC method has been developed by integrating preprocessing, feature selection, and classification.
- To minimize air pollution prediction time, the MNOLPELMC method integrates data preprocessing, significant feature selection into the hidden layer of the extreme learning machine classifier model.
- To improve the accuracy and minimize root mean square error, the Otsuka–Ochiai similarity coefficient is employed for evaluating training and testing data samples, providing multiple classification outcomes at the output layer.

2. Related Works

A Support Vector Machine (SVM) model was developed in [16]. However, it failed to apply the ensemble machine learning classification algorithm for analyzing the various meteorological parameters, like temperature and wind speed, when predicting the AQI class.

Machine learning methods were developed in [17]. However, it failed to forecast the pollutant concentrations for decision-making based on expected air quality conditions. The Long Short-Term Memory (LSTM) model was developed in [18]. However, the feature selection process remained unaddressed for further improving

prediction accuracy. An integration of Spring Search Optimization with a Hybrid Deep Learning (BSSO-HDL) model was designed in [19]. However, it failed to incorporate large volumes of data for prediction. Hybrid deep learning models were developed in [20].

However, the designed models exhibited the slowest rate of performance degradation during multi-step prediction. A hybrid deep learning and quantum-inspired particle swarm optimization model was developed in [21]. However, it failed to focus on handling the long-sequence data for temporal patterns and dependencies analysis.

A multi-modal framework was developed in [22] using a deep learning model. However, it failed to investigate the extent of performance improvement from the insertion of features from weather stations. A deep learning model was designed in [23].

However, it failed to integrate external variables such as weather patterns and industrial operations to further improve the accuracy of predictions. A Machine Learning (ML)-based Artificial Neural Network (ANN) was developed in [24] for IoT-based intelligent pollution monitoring. However, time-efficient pollution forecasting remained a major concern. Machine learning techniques were utilized in [25].

However, deep learning and ensemble methods were not considered for predicting other pollutants. Regression models were constructed in [26]. But it failed to improve the accuracy of the approach with minimal hardware support.

A two-stage attention mechanism was introduced in [27] based on transfer learning. However, it failed to consider the generalization and robustness of the model to further improve the overall performance of the prediction model. Advanced tree-based machine learning models were developed in [28].

However, the employed models failed to consider various input parameters for accurately predicting air pollutant levels. Advanced deep learning models were introduced in [29]. However, time complexity analysis in PM2.5 concentration prediction remained a major challenge. A new neural network structure model was developed in [30]. However, it failed to consider air temperature, solar radiation, and gas usage for prediction.

3. Proposal Methodology

Air pollution is the occurrence of substances in the atmosphere, and it poses considerable risks to human health and the environment. The proposed MNOLPELMC method aims to provide an accurate prediction of air pollution in a more efficient and cost-effective manner with minimal time consumption.

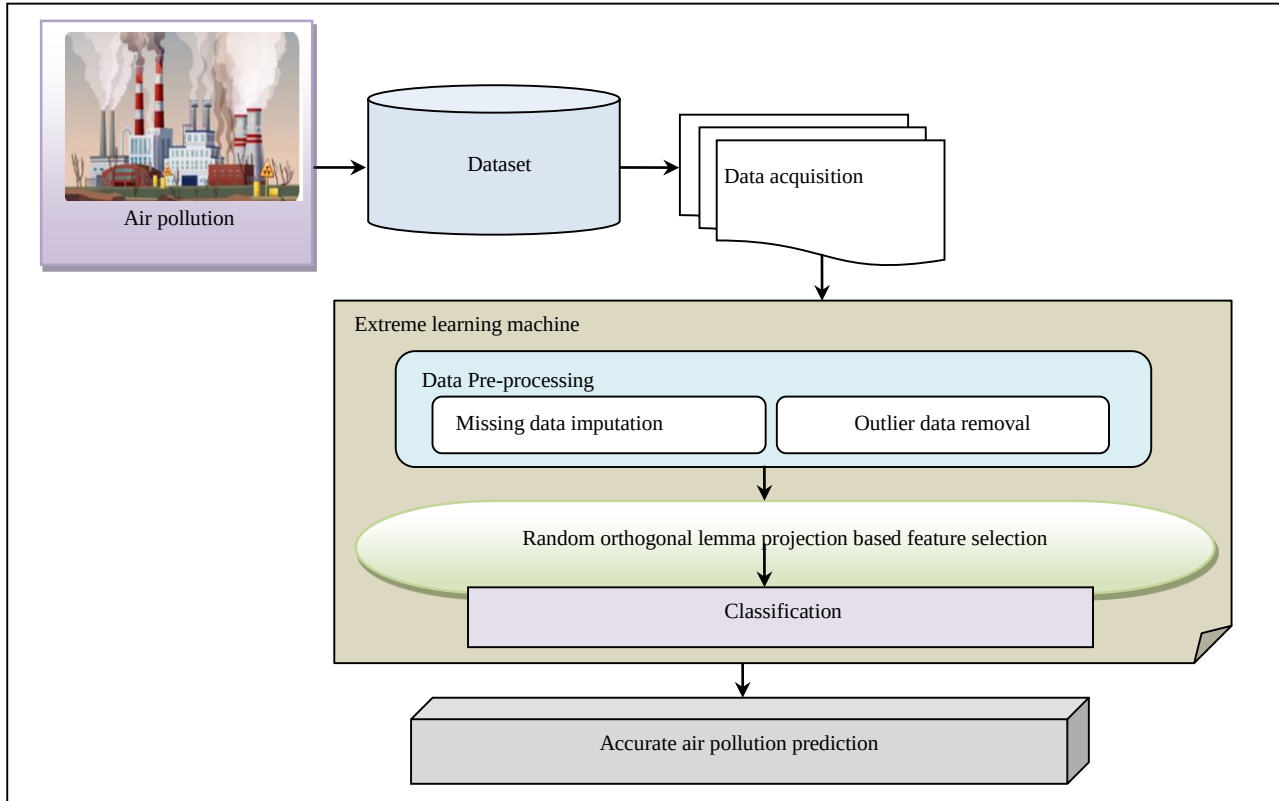


Fig. 1 Overview of the architecture displaying the steps involved in the MNOLPELMC method

In Figure 1, the MNOLPELMC method is divided as data acquisition, preprocessing, feature extraction, and classification.

These phases work to predict air pollution as illustrated in Figure 1. The MNOLPELMC method focusing on its ability to improve the accuracy of air pollution prediction with minimal time complexity.

3.1. Data Acquisition

Data acquisition is the fundamental step in the proposed MNOLPELMC method that involves collecting a large volume of data samples from the repository. In the proposed model, the Global Air Pollution Dataset is used, and it is taken from <https://www.kaggle.com/datasets/hasibalmuzdadid/global-air-pollution-dataset>.

Table 1. Attribute information

S. No	Features	Description
1	Country	Name of the country
2	City	Name of the city
3	AQI Value	Overall AQI value of the city
4	AQI Category	Overall AQI category of the city
5	CO AQI Value	AQI value of Carbon Monoxide in the city
6	CO AQI Category	QI category of Carbon Monoxide in the city
7	Ozone AQI Value	AQI category of Ozone of the city
8	Ozone AQI Category	AQI category of Ozone of the city
9	NO2 AQI Value	AQI value of Nitrogen Dioxide in the city
10	NO2 AQI Category	AQI category of Nitrogen Dioxide in the city
11	PM2.5 AQI Value	The QI value of Particulate Matter with a diameter of 2.5 micrometers or less in the city
12	PM2.5 AQI Category	QI category of Particulate Matter with a diameter of 2.5 micrometers or less in the city

The input dataset 'DT' is prepared in the form of a matrix expressed as follows,

$$IM = \begin{bmatrix} A_1 & A_2 & \dots & A_m \\ DS_{11} & DS_{12} & \dots & DS_{1n} \\ DS_{21} & DS_{22} & \dots & DS_{2n} \\ \vdots & \vdots & \dots & \vdots \\ DS_{m1} & DS_{m2} & \dots & DS_{mn} \end{bmatrix} \quad (1)$$

From (1), the input matrix 'IM' is structured based on 'm' number of features. $A_1, A_2, \dots, A_m$ and the column represents overall sample instances, data, or records ' DS_n ' presented in

the 'n' rows respectively. The MNOLPELMC method utilizes the extreme learning classifier with five different layers, namely one input layer, one output layer, and three hidden layers. The extreme learning model is a feed-forward neural network having multiple layers of hidden nodes for the classification of the data samples. The Extreme Learning Machine algorithm (ELM) helps to provide better performance at an extremely fast learning speed than the other deep learning models. Therefore, the proposed method utilizes the MNOLPELMC method for air pollution prediction with minimal time consumption.

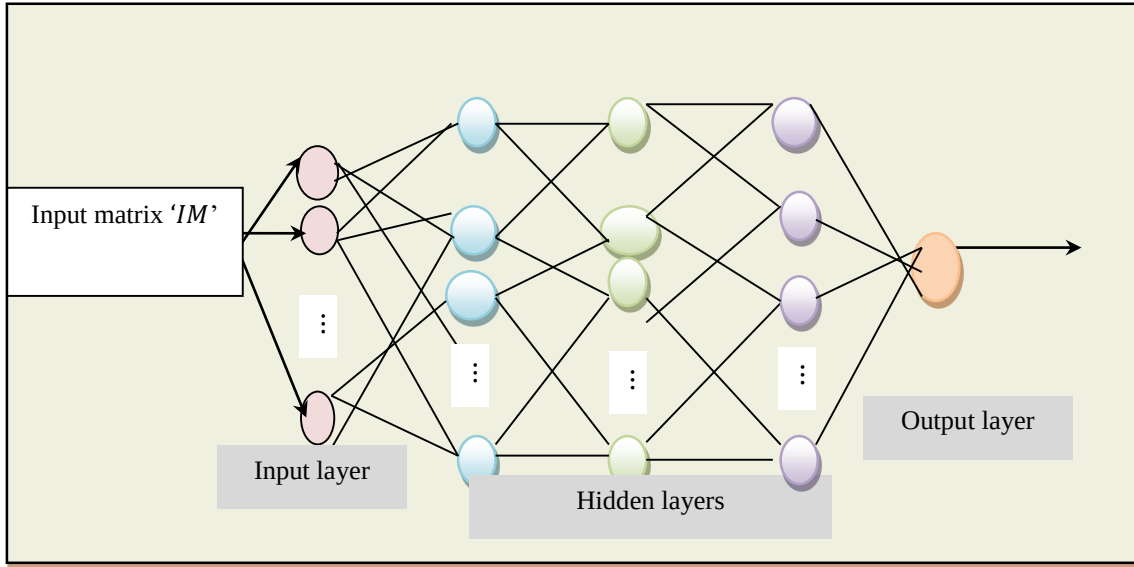


Fig. 2 Schematic construction of the extreme learning machine algorithm

Figure 2 illustrates the schematic construction of an extreme learning machine classifier, which consists of an input layer, one or more hidden (middle) layers, and an output layer. The layer consists of a neuron to transfer the data between the layers.

Each neuron in the layers is fully connected to other layers by means of a hyperparameter called weights. The neuron in the input layer received the input data samples.

Three hidden layers are used in this structure, as shown in Figure 2 for preprocessing, feature selection, and data classification, respectively. Finally, the output layer provides the final prediction outcomes. The extreme learning machine network considers the training set. $\{DS_i, Y_j\}$ where DS indicates a training data sample's $DS_i = \{DS_1, DS_2, \dots, DS_n\}$ and Y_j refers to the final air pollution prediction outcomes. Each neuron in the input layer receives training data samples and forwards them to the hidden layer, where the hyperparameters, such as weights and biases, are assigned. Neurons in the hidden layer compute a weighted sum of the input data samples received from the input layer. It is mathematically expressed as,

$$Z = \sum_{i=1}^n (DS_i * W_{ij}) + b \quad (2)$$

Where Z represents a weighted sum output, DS represents an input data sample, W_{ij} indicates a weight between neuron i in the input layer and neuron j in the hidden layer, and b represents the bias. In the first hidden layer, the proposed MNOLPELMC method performs the data preprocessing that involves organizing the raw data samples for further analysis. The main aim of data preprocessing is to clean and transform the data into a suitable format for providing accurate and significant results. Preprocessing helps to enhance the quality and efficiency of the dataset and minimizes the likelihood of errors. The Preprocessing step includes two major processes, namely missing data handling and outlier data detection. Missing data is referred to as an absence of data or a null value for a particular feature within a dataset. The missing data handling process is significant in the context of prediction since it generates inaccurate results. The proposed method utilizes the Sibson interpolation method for handling the missing data. Sibson interpolation, also called as proximity interpolation, provides a smoother approximation to the underlying accurate function. Therefore, the function of Sibson interpolation is expressed as

$$S = \sum_{i=1}^k Ds_i \delta_i \quad (3)$$

Where S represents a Sibson interpolation for handling the missing data, δ_i indicates a weight assigned to neighboring data points ' Ds_i ', k designates a number of neighboring data points. The weights are assigned based on the similarity measure between the data samples. The similarity is measured as follows.

$$Sim = \frac{[Ds_j \cap Ds_k]}{\sum Ds_j + \sum Ds_k - [Ds_j \cap Ds_k]} \quad (4)$$

Where, ' Sim ' indicates a Jaccard index coefficient, Ds_j and Ds_k denotes the neighboring data samples, $Ds_j \cap Ds_k$ denotes a mutual dependence between the two samples. The coefficient (Sim) returns the output value from 0 to 1. Higher weight is assigned to higher similarity, while lower weight is given to less similarity. Followed by, Grubbs maximum normalized residual test is employed for outlier data removal. Outlier data is a data point that is considerably dissimilar from the rest of the data within the dataset. Outliers are caused the errors during the prediction process. Therefore, these data points are removed before performing any computational tasks. Grubbs maximum normalized residual test is a statistical analysis used to distinguish the outliers within the dataset. The statistical analysis is expressed as follows,

$$R = \arg \max \left[\frac{Ds_i - \mu_{Ds}}{\sigma} \right] \quad (5)$$

Where ' R ' denotes a statistical outcome, Ds_i indicates a data sample, μ_{Ds} denotes a mean value of the data samples, σ indicates a deviation between the data samples and the mean value.

From the analysis, the maximum statistical results of the particular data samples are called as an outlier, and it is removed from the dataset. Therefore, dataset preprocessing helps to improve the accuracy and quality of a dataset.

3.2. Feature Selection

The preprocessed results are transferred into the second hidden layer of the extreme learning machine classifier model for selecting the significant features from the dataset. The feature selection process of the proposed model helps to reduce the dimensionality of the dataset by removing the irrelevant features. This process helps to reduce the time complexity of the air pollution prediction.

The proposed method utilizes the random orthogonal lemma projection process in the hidden layer 2 for feature selection. The lemma projection is the process of projecting the set of data points from the high-dimensional space into lower dimension. A simple projection strategy has some noise bounds due to the loss of information. On the contrary, orthogonal projection is more applicable to raw data samples

and latent feature representations. The proposed method considers the features and their data samples are in matrix format. After that, a random matrix is constructed with a Gaussian distribution function as,

$$F(Ds|\mu, \sigma) = \frac{1}{\sigma \sqrt{2\pi}} \exp \left[-\frac{|Ds_i - \mu|^2}{\sigma^2} \right] \quad (6)$$

Where, $F(Ds|\mu, \sigma)$ denotes a Gaussian distribution function with the mean ' $\mu = 0$ ' and deviation ' $\sigma = 1$ '. Based on the above distribution measure, the random matrix is constructed as follows,

$$R = \begin{bmatrix} F_{11} & F_{12} & \dots & F_{1n} \\ F_{21} & F_{22} & \dots & F_{2n} \\ \vdots & \vdots & \dots & \vdots \\ F_{m1} & F_{m2} & \dots & F_{mn} \end{bmatrix} \quad (7)$$

Where R indicates a random matrix. Followed by, the orthogonal matrix is generated as follows,

$$Q = R^T R \quad (8)$$

Where Q indicates an orthogonal random matrix, R^T indicates a transpose of matrix ' R '. In order to perform data projection, the generated orthogonal random matrix is multiplied with input data matrix.

$$P = Q * A \quad (9)$$

Where P denotes a projected new matrix in low-dimensional space, Q indicates an orthogonal random matrix, A denotes an input feature matrix. After that, the variance between the feature values is measured in the low-dimensional space. The proposed model utilizes the Wilcoxon index function for computing the variance between the values. The Wilcoxon index is a mathematical function used to measure the deviation between the two samples.

$$WI = \sum_{i=1}^n \sum_{j=1}^m |A_i - A_j| \quad (10)$$

Where WI denotes a Wilcoxon index function that provides the variance results between 0 and 1, finally, the high-variance features are selected as significant features for air pollution prediction. The relevant feature selection process with high variance is selected as follows,

$$Z = \begin{cases} WI > Th & ; \text{features selected} \\ WI < Th & ; \text{not selected} \end{cases} \quad (11)$$

From (11), the output of the random orthogonal lemma projection ' Z ' is obtained by setting the threshold value for the variance ' Th '. Therefore, the variance between the features ' WI ' is greater than the threshold is used for selecting the relevant features.

Otherwise, the features with low variances are not selected, and it removed from the dataset. From the Table 1,

selected significant features after the projection are listed in Table 2.

Table 2. Selected feature list

S. No	Features	Description
1	City	Name of the city
2	AQI Value	Overall AQI value of the city
3	CO AQI Value	AQI value of Carbon Monoxide in the city
4	Ozone AQI Value	AQI value of Ozone of the city
5	NO2 AQI Value	AQI value of Nitrogen Dioxide in the city
6	PM2.5 AQI Value	The QI value of Particulate Matter with a diameter of 2.5 micrometers or less in the city
7	PM2.5 AQI Category	QI category of Particulate Matter with a diameter of 2.5 micrometers or less in the city

3.3. Classification

After selecting significant features, the classification process is carried out for predicting the air pollution level. The classification process is done by using training data samples and testing data samples. The proposed method utilizes the Otsuka–Ochiai similarity for analyzing the relationship between the testing and training data samples.

$$\rho = \frac{\sum Ds_{tr}Ds_{ts}}{\sqrt{\sum Ds_{tr}^2} \sqrt{\sum Ds_{ts}^2}} \quad (12)$$

Where ρ indicates the Otsuka–Ochiai similarity coefficient, Ds_{tr} symbolizes the training data samples and Ds_{ts} denotes a testing data samples, $\sum Ds_{tr}Ds_{ts}$ denotes the sum of the product of the paired score of Ds_{tr} and Ds_{ts} , $\sum Ds_{tr}^2$ denotes the sum of the squared scores of the Ds_{tr} , $\sum Ds_{ts}^2$ indicates a sum of the squared scores of the Ds_{ts} . The Orchini similarity coefficient (ρ) provides the output value between 0 and 1. Here, 1 indicates higher similarity between the two data samples, and 0 represents less similarity.

Based on the analysis, the data samples are classified as good, moderate, unhealthy, and very unhealthy. The good indicates the air is clean and poses minimal health risks to most individuals. The moderate results indicate acceptable air pollution with potential health concerns. An Unhealthy and very unhealthy AQI level indicates a significant weakening in air quality, where the high pollution initiates to affect the health of the individuals. Therefore, accurate outcomes of the air pollution prediction results, good, moderate, unhealthy, and very unhealthy, are obtained at the output layer with softmax activation functions.

$$f_{Soft_Max} = \frac{\exp(Y_k)}{\sum_{j=1}^k \exp(Y_j)} \quad (13)$$

Where, softmax activation function ' f_{Soft_Max} ' softmax activation function at the output layer is used to construct multiple classification outcomes, Y_k indicates an output for the j^{th} class, k denotes the total number of classes used for air pollution prediction. In this way, accurate air pollution prediction results are obtained at the output layer.

Algorithm 1 Illustrates different processes used for predicting various types of air pollution levels

// Algorithm 1: Maximum Normalized Orthogonal Lemma Projective Extreme Learning Machine Classification Method
Input: Dataset 'D', features $A_1, A_2, A_3, \dots, A_m$, data samples or instances $Ds_1, Ds_2, Ds_3, \dots, Ds_n$
Output: Increase the air pollution prediction accuracy
Begin
Step 1: Collect number of features ' $A_1, A_2, \dots, A_m$ ' and samples ' $Ds_1, Ds_2, \dots, Ds_n$ '- input layer
Step 2: For each sample Ds_i
Step 3: Compute the weighted sum by the neuron using (1)
Step 4: End For
Step 5: if any missing data is found, then-----[hidden layer 1]
Step 6: ApplySibson interpolation using (3)
Step 7: end if
Step 8: Measure Grubbs maximum normalized residual test using (5)
Step 9: Find outliers and removed
Step 10: For each preprocessed output -----[hidden layer 2]
Step 11: Create random matrix with a Gaussian distribution function using (7)
Step 12: Create the orthogonal random matrix (8)
Step 13: Obtain the projected matrix (9)
Step 14: Compute the Wilcox index variance using (10)
Step 15: if ($WI > Th$)then

Step 16: Features are selected as relevant
 Step 17: else
 Step 18: Features are not selected as relevant
 Step 19: End if
 Step 20: For each training and testing data samples -----[hidden layer 3]
 Step 21: Measure the similarity using (12)
 Step 22: Obtain the multiple classification results
 Step 23: Obtain final prediction results with softmax activation function at the output layer using (13)
 End

4. Experimental Settings

Experimental assessment of the MNOLPELMC method, along with existing methods NF-VAE [1], Neural transfer learning model [2], is implemented in Python. using the Global Air Pollution Dataset taken from <https://www.kaggle.com/datasets/hasibalmuzdadid/global-air-pollution-dataset>.

5. Performance Analysis

Performance analysis of the MNOLPELMC method in comparison to existing methods [1, 2]. The evaluation is conducted using various metrics.

Air pollution prediction accuracy: It is measured as the ratio of data samples for which air pollutions levels are accurately predicted to the total number of data collected from the dataset. It is expressed as follows,

$$APPA = \sum_{i=1}^n \left(\frac{DCP}{Ds_i} \right) * 100 \quad (14)$$

Precision: It is an efficient metric used for computing the accuracy of model positive predictions. It is calculated as the ratio of True Positives (TP) to the total number of predicted positive samples, which includes both True Positives (TP) as well as False Positives (FP).

$$PS = \frac{TP}{TP+FP} \quad (15)$$

Recall: it is also known as sensitivity, which evaluates model efficiency for accurately identifying positive instances. The formula for computing the recall is given below,

$$RL = \frac{TP}{TP+FN} \quad (16)$$

Root Mean Squared Error: It is measured as the ratio of the number of data samples where the difference between actual results matches the predicted results to the total number of data samples. It is computed as,

$$RMSE = \left[\sqrt{\frac{(Y_{Act} - Y_{Pre})^2}{n}} \right] \quad (17)$$

Prediction time: It is measured as the amount of time required for predicting the air pollution based on input data samples. The overall prediction time is computed as follows,

$$PT = \sum_{i=1}^n Ds_i * Time [APP] \quad (18)$$

Table 3 depicts the analysis of the air pollution prediction accuracy of three methods. The average of the seten comparison results illustrates prediction accuracy of MNOLPELMC improved by 3% and 5% than the [1, 2].

Table 3. Comparison of air pollution prediction accuracy versus the number of data samples

Number of data samples	Air pollution prediction accuracy (%)		
	MNOLPELMC	GNF-VAE	Neural transfer learning model
2000	93.1	91.2	90.5
4000	93.85	90.32	88.65
6000	94.11	91.05	89.03
8000	93.75	91.98	88.73
10000	94.36	91.65	89.23
12000	93.40	90.05	88.05
14000	94.47	91.05	88.65
16000	93.55	90.73	88.13
18000	92.31	90.77	88.11
20000	93.05	91.94	89.78

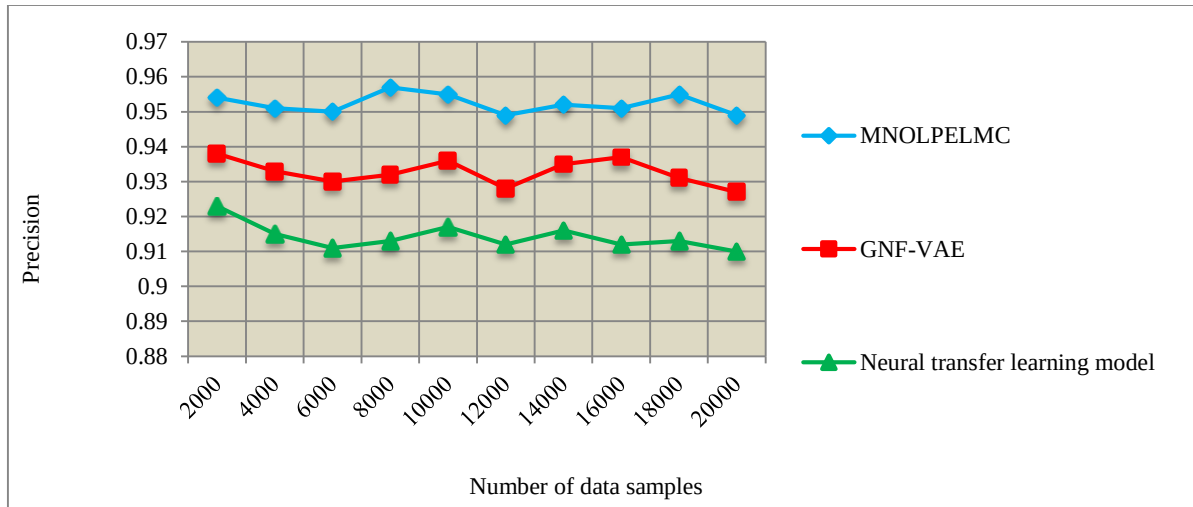


Fig. 3 Performance analyses of precision

Figure 3 portrays graphical outcomes of precision in predicting air pollution level using three methods. A

comparison of these results shows that the MNOLPELMC method increases precision by 2% and 4% than [1, 2].

Table 4. Comparison of recall versus number of data samples

Number of data samples	Recall		
	MNOLPELMC	GNF-VAE	Neural transfer learning model
2000	0.968	0.953	0.937
4000	0.96	0.951	0.932
6000	0.962	0.945	0.925
8000	0.963	0.949	0.928
10000	0.958	0.935	0.917
12000	0.965	0.942	0.922
14000	0.955	0.938	0.916
16000	0.963	0.947	0.929
18000	0.968	0.945	0.926
20000	0.965	0.942	0.925

Table 4 illustrates the performance outcomes of recall with number of data samples ranges. The comparison results

shows performance of recall using the MNOLPELMC method is improved by 2% and 4% than the [1, 2].

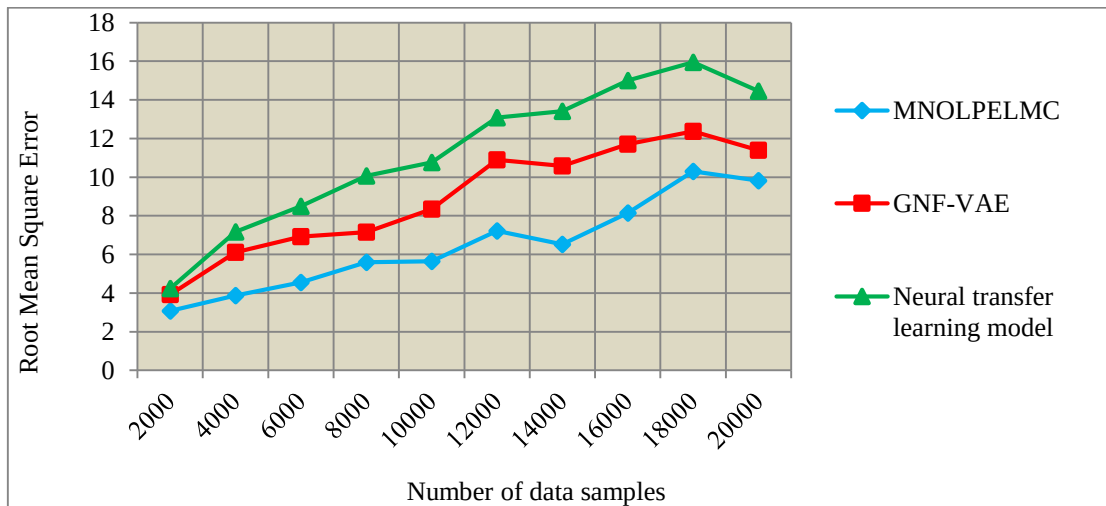


Fig. 4 Performance analyses of root mean squared error

Figure 4 illustrates experimental results of RMSE versus the number of data samples. Overall, comparable results reveal MNOLPELMC minimizes RMSE in accurately predicting air pollution by 28% and 42% than [1, 2].

Table 5 depicts performance outcomes of air pollution prediction time versus the number of data. The average of ten results reveals that MNOLPELMC minimizes overall prediction time by 10% and 16% than the [1, 2].

Table 5. Comparison of prediction time versus number of data samples

Number of data samples	Prediction time (ms)		
	MNOLPELMC	GNF-VAE	Neural transfer learning model
2000	28	32	34
4000	32.5	35.6	38.9
6000	36.5	40.2	42.6
8000	40.2	45.6	48.5
10000	43.6	48.5	52.3
12000	45.8	50.3	55.8
14000	48.3	52.5	57.2
16000	50.2	55.6	58.3
18000	55.6	60.7	63.5
20000	58.4	63.5	68.6

6. Conclusion

Accurate detection and forecasting of air pollution concentration levels play a vital factor in protecting public health as well as the environment. The MNOLPELMC method is designed to improve the accuracy of air pollution predictions while reducing time consumption.

The results exhibit that the MNOLPELMC method outperforms traditional deep learning models, showing notable improvements in accuracy, precision, and recall. The MNOLPELMC method reduces prediction time and RMSE compared to conventional approaches.

References

- [1] Prasanjit Dey, Soumyabrata Dev, and Bianca Schoen Phelan, "Predicting Multivariate Air Pollution: A Gaussian-Mixture Nested Factorial Variational Autoencoder Approach," *IEEE Geoscience and Remote Sensing Letters*, vol. 21, pp. 1-5, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [2] Idriss Jairi et al., "Enhancing Air Pollution Prediction: A Neural Transfer Learning Approach Across Different Air Pollutants," *Environmental Technology and Innovation*, vol. 36, pp. 1-22, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [3] Kuldeep Singh Rautela, and Manish Kumar Goyal, "Transforming Air Pollution Management in India with AI and Machine Learning Technologies," *Scientific Reports*, vol. 14, no. 1, pp. 1-19, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [4] S. Abu El-Magd et al., "Environmental Hazard Assessment and Monitoring for Air Pollution using Machine Learning and Remote Sensing," *International Journal of Environmental Science and Technology*, vol. 20, no. 6, pp. 6103-6116, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [5] K. Kumar, and B.P. Pande, "Air Pollution Prediction with Machine Learning: A Case Study of Indian Cities," *International Journal of Environmental Science and Technology*, vol. 20, no. 5, pp. 5333-5348, 2022. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [6] Vasileios N. Matthaios et al., "Predicting Real-Time Within-Vehicle Air Pollution Exposure with Mass-Balance and Machine Learning Approaches using On-Road and Air Quality Data," *Atmospheric Environment*, vol. 318, pp. 1-12, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [7] Tania Septi Anggraini et al., "Machine Learning-based Global Air Quality Index Development using Remote Sensing and Ground-based Stations," *Environmental Advances*, vol. 15, pp. 1-13, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [8] Manzoor Ansari, and Mansaf Alam, "An Intelligent IoT-Cloud-based Air Pollution Forecasting Model using Univariate Time-Series Analysis," *Arabian Journal for Science and Engineering*, vol. 49, no. 3, pp. 3135-3162, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [9] Liam Jordan Berrisford, Hugo Barbosa, and Ronaldo Menezes, "A Data-Driven Supervised Machine Learning Approach to Estimating Global Ambient Air Pollution Concentrations with Associated Prediction Intervals," *Royal Society Open Science*, vol. 12, no. 7, pp. 1-33, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [10] Farzaneh Mohammadi et al., "Prediction of Atmospheric PM_{2.5} level by Machine Learning Techniques in Isfahan, Iran," *Scientific Reports*, vol. 14, no. 1, pp. 1-12, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [11] Kahraman Oğuz, and Muhammet Ali Pekin, "Prediction of Air Pollution with Machine Learning Algorithms," *Turkish Journal of Science and Technology*, vol. 19, no. 1, pp. 1-12, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]

- [12] Ignacio-Iker Prado-Rujas et al., “A Multivariable Sensor-Agnostic Framework for Spatio-Temporal Air Quality Forecasting based on Deep Learning,” *Engineering Applications of Artificial Intelligence*, vol. 127, pp. 1-17, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [13] Bainomugisha et al., “AI-Driven Environmental Sensor Networks and Digital Platforms for Urban Air Pollution Monitoring and Modeling,” *Societal Impacts*, vol. 3, pp. 1-6, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [14] M.J. Jiménez-Navarro et al., “Explaining Deep Learning Models for Ozone Pollution Prediction Via Embedded Feature Selection,” *Applied Soft Computing*, vol. 157, pp. 1-11, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [15] Qutaiba Humadi Mohamme, and Anupama Namburu, “Hybrid Model and Framework for Predicting Air Pollutants in Smart Cities,” *Journal of Engineering and Sustainable Development*, vol. 28, no. 3, pp. 392-406, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [16] Mohsin Imam et al., “Air Quality Monitoring using Statistical Learning Models for Sustainable Environment,” *Intelligent Systems with Applications*, vol. 22, pp. 1-15, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [17] A. Samad et al., “Air Pollution Prediction using Machine Learning Techniques – An Approach to Replace Existing Monitoring Stations with Virtual Monitoring Stations,” *Atmospheric Environment*, vol. 310, pp. 1-18, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [18] G. Naresh, and B. Indira, “Air Pollution Prediction using Multivariate LSTM Deep Learning Model,” *International Journal of Intelligent Systems and Applications in Engineering*, vol. 12, no. 8s, pp. 211-220, 2024. [[Google Scholar](#)] [[Publisher Link](#)]
- [19] Sreenivasulu Kutala et al., “Hybrid Deep Learning-based Air Pollution Prediction and Index Classification using an Optimization Algorithm,” *AIMS Environmental Science*, vol. 11, no. 4, pp. 551-575, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [20] Gang Chen et al., “A Hybrid Deep Learning Air Pollution Prediction Approach based on Neighborhood Selection and Spatio-Temporal Attention,” *Scientific Reports*, vol. 15, no. 1, pp. 1-20, 2025. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [21] Anh Tuan Nguyen et al., “Predicting Air Quality Index using Attention Hybrid Deep Learning and Quantum-Inspired Particle Swarm Optimization,” *Journal of Big Data*, vol. 11, no. 1, pp. 1-38, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [22] Saad Hameed et al., “Deep Learning based Multimodal Urban Air Quality Prediction and Traffic Analytics,” *Scientific Reports*, vol. 13, no. 1, pp. 1-19, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [23] Qing Tao et al., “Air Pollution Forecasting using a Deep Learning Model based on 1D Convnets and Bidirectional GRU,” *IEEE Access*, vol. 7, pp. 76690-76698, 2019. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [24] Kianat, Humayun Salahuddin, and Muhammad Saleem Anjum, “IoT based Intelligent Pollution Monitoring System using Machine Learning Technique,” *Journal of Computing and Biomedical Informatics*, vol. 6, no. 2, pp. 529-537, 2024. [[Google Scholar](#)] [[Publisher Link](#)]
- [25] Tshepang Duncan Morapedi, and Ibidun Christiana Obagbuwa, “Air Pollution Particulate Matter (PM2.5) Prediction in South African Cities using Machine Learning Techniques,” *Frontiers Artificial Intelligence*, vol. 6, pp. 1-19, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [26] S. John Livingston et al., “An Ensembled Method for Air Quality Monitoring and Control using Machine Learning,” *Measurement: Sensors*, vol. 30, pp. 1-6, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [27] Zhanfei Ma et al., “Air Pollutant Prediction Model based on Transfer Learning Two-Stage Attention Mechanism,” *Scientific Reports*, vol. 14, no. 1, pp. 1-18, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [28] Caroline Matara et al., “Prediction of Vehicle-induced Air Pollution based on Advanced Machine Learning Models,” *Engineering, Technology and Applied Science Research*, vol. 14, no. 1, pp. 12837-12843, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [29] Zihang Gao, Xinyue Mo, and Huan Li, “Prediction of PM_{2.5} Concentration based on Deep Learning, Multi-Objective Optimization, and Ensemble Forecast,” *Sustainability*, vol. 16, no. 11, pp. 1-15, 2024. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]
- [30] Wenlin Li, and Xuchu Jiang, “Prediction of Air Pollutant Concentrations based on TCN-BiLSTM-DMAAttention with STL Decomposition,” *Scientific Reports*, vol. 13, no. 1, pp. 1-12, 2023. [[CrossRef](#)] [[Google Scholar](#)] [[Publisher Link](#)]