



Blog Feedback Prediction based on Ensemble Machine Learning Regression Model: Towards Data Fusion Analysis

Hamzah A. Alsayadi^{1,*}, El-Sayed M. El-Kenawy², Abdelhameed Ibrahim³, Marwa M. Eid⁴, Abdelaziz A. Abdelhamid⁵

¹Computer Science Department, Faculty of Sciences, Ibb University, Yemen

²Department of Communications and Electronics, Delta Higher Institute of Engineering and Technology, Mansoura, 35111, Egypt

³Computer Engineering and Control Systems Department, Faculty of Engineering, Mansoura University, 35516, Mansoura Egypt

⁴Faculty of Artificial Intelligence, Delta University for Science and Technology, Mansoura, Egypt

⁵Computer Science Department, Faculty of Computer and Information Sciences, Ain Shams University, Cairo, 11566, Egypt

Emails: hamzah.sayadi@cis.asu.edu.eg; skenawy@ieee.org; afai79@mans.edu.eg; marwa.3eed@gmail.com; abdelaziz@cis.asu.edu.eg

Abstract

The last decade lead to an unbelievable growth of the importance of social media. Due to the huge amounts of documents appearing in social media, there is an enormous need for the automatic analysis of such documents. In this work, we proposed various regression models for the blog feedback prediction to be used in the data fusion environment. These models include decision tree regressor, MLP regressor, SVR, random forest regressor, and K-Neighbors regressor. The models are enhanced by average ensemble and ensemble using K-Neighbors regressor. The Blog Feedback dataset is used for training and evaluating the proposed models. The results show that there is a decrease in RMSE, MAE, MBE, R, R2, RRMSE, NSE, and WI when compared to the traditional methods.

Keywords: Blog feedback prediction; Ensemble model; Machine learning; Regression model; Data fusion

1 Introduction

The significance of social media has increased incomprehensibly during the past ten years. While early social media platforms like blogs, tweets, Facebook, YouTube, and social tagging systems mostly served as a form of entertainment for a small number of ardent users, news that circulates on these platforms today may influence the most significant societal changes, like the Islamic world's revolutions or the US presidential elections. Additionally, through social media platforms, advertising and news about new goods, services, and businesses are disseminated swiftly. On the one hand, this may be a fantastic opportunity to market fresh goods and services [1, 2, 3].

Additionally, sociological studies have shown that negative perceptions travel far more fast than good ones, thus if

negative news about a firm appears in social media, the corporation may need to act immediately to prevent losses [1]. However, it is now a goldmine for fostering relationships among users and pursuing their interests. As a result, our communication, sharing of experiences, and opinions have changed, and it has shown to be a valuable resource for government, marketing, and educational institutions [4, 5].

Feedback of blogs is collected in the form of likes, dislikes, comments, shares, and clicks. People's opinions depend on the topic, the substance, the time(s), the type, the interest, and the number of followers connected to that. Additionally, the profile of the bloggers, their posts, and their networks are related to the comments from blogs [6]. The use of a blog as an example in social media and potential determinants of its popularity, such as likes and comments, etc.

The analysis of human behavior, which is a key component of blog prediction, helps to foster peoples' interest in particular subjects. This concept could improve how blogs' readership trends are examined. Even new bloggers can evaluate the behaviors of their followers and select the most important concepts to go viral in a short amount of time [2]. The number of comments acts as a stand-in for attention in the blog feedback prediction issue [1, 3], where the goal is to forecast the number of comments that a blog will receive. Different features that relate to the textual content, source (blog or website), or other circumstances, such as the day of the week the post was published, are utilized to predict the number of comments [7].

Because there are so many documents on social media, it is impossible for human professionals to analyze them all, hence there is a great need for robotic analysis of these materials. However, there are some unique characteristics of the application area that must be taken into consideration for the study. Particularly, the unstructured, dynamic, and quickly-changing content of social media documents, such as when a blog item publishes and users can comment on it right away [1]. Recently, the blog feedback prediction task received considerable attention, see e.g. [8, 9, 10]. Various models have been used, such as state-of-the-art variants of fuzzy systems [8], neural networks, and decision trees [10].

In this work, the blog feedback prediction has been done using different regression methods. We proposed five different models that include regression models such as decision tree regressor, multi-Layer perceptron Model (MLP) regressor, support vector regression (SVR), random forest regressor, and K-Neighbors regressor. In addition, the ensemble model is proposed to optimize the parameters of the regression models. Also, the proposed have been evaluated using statistical metrics such as Mean Absolute Error (MAE), Mean Absolute Percentage Error (MAPE), Mean Square Error (MSE), Root Mean Square Error (RMSE), and R². Results show the achievement of better performance with decreased error rate when compared to traditional prediction models.

The rest of the paper is organized as follows. In Section 2, we present the literature review. In Section 3, the methodology is presented. Section 4 includes the experimental setup. Finally, Section 5 includes the conclusion and suggestions.

2 Literature Review

For many years, various studies have been done related to blog feedback prediction. In [1], Krisztian Buza forecasted how many comments a blog post will receive. This study is the most closely related to our own. The goal is to model samples from recently uploaded blog posts to forecast how many responses a blog post will receive over the course of the following H hours. To identify the four traits, the articles have been extracted. Seven predicted models were utilized, including (i) neural networks, (ii) RBF networks, (iii) regression trees, (iv) closest neighbours, (v) multivariate linear regression, (vi) bagging out of the ensemble, and (vii) SVM. The fundamental, textual, weekday, and parent features were also taken into consideration.

By combining real-time, SenTime, and batch data processing for opinion mining, the authors of [11] presented the identification and classification of user behaviour based on Facebook comments as a data fusion approach. The 2016 US Presidential Elections were the chosen topic. 200 posts were gathered from the Facebook pages of the BBC and

CNN, two renowned sources. The findings show how attitudes regarding a subject vary over time, and they represent three different sorts of opinions throughout a real-time stream. Results are assessed using the Mean Absolute Error (MAE) of the forecast.

According to user replies and reaching trends, the authors of [12] proposed predicting the lifespan of internet news stories. It demonstrates how accurately forecasting future replies to news stories may be supported by identifying the first 10 to 20 minutes of responses. The number of responses a news piece will receive after it is published has been predicted using linear regression models.

Authors in [13] predicted box office using microblogs. First, the box office dataset was extracted from the Entgroup box office website, while the microblog data was extracted from Tencent Microblog. Positive, negative, and neutral comments have been separated into three groups. SVM and neural networks have been used to analyse the data, and their performance has been assessed using correlation and average MAPE.

A popular social networking service, Facebook, has been predicted to receive a high volume of answers, according to authors in [10]. In contrast to other models like the Multi-Layer Perceptron, RBF networks, M5P and REP trees, they have used data for forecasting from Hadoop clusters and forecasted using regression models, Neural Networks (NN), and Decision Trees (DT). They separated it into four feature sets, including Page, Essential, Weekday, and Basic features, after performing four data pre-processing modules. Hits@10, AUC@10, Mean Absolute Error (MAE), and Time were employed as evaluation criteria in this study (s). According to their analysis, decision trees perform better than alternative regression models across the board.

The paper in [14] talks about categorizing blog comment responses in two categories. Neural networks and Adaboost algorithms are used for classification on a benchmark feedback blog dataset that is publicly available online. Based on the accuracy of both algorithms' prediction rates, it has been thought of extracting the four feature sets from the dataset. The computational costs of both techniques have also been assessed for all feature sets. Results of the experiment demonstrate that, with a prediction rate of 91.4%, the Adaboost classifier performs better than neural networks in all respects.

The study in [2] proposed using the UCI BlogFeedback dataset and an Adaptive Neuro Fuzzy Inference System to estimate the amount of feedback blogs. This proof-of-prediction concept's accuracy for Accuracy, Correlation, and Mean Squared Error has been compared with ANN, SVM, and Random Forest. It can attain 94% accuracy, which makes it an excellent choice for industrial social media analysis applications.

The blog feedback prediction challenge was taken into account, and factorization machines for this purpose were extensively studied in study [3]. Factorization machines are effective models for predicting blog comments, as one might assume. On the other hand, our investigation showed that the development of effective training algorithms may be difficult, as it seems that interaction weights may only be learned effectively if feature weights already have acceptable values.

The tSVR algorithm of local SVR was introduced by the authors in [15] and gets excellent results for non-linear regression of big datasets. To train the tSVR, the entire training dataset must be divided into k terminal-nodes. In training local SVR, this seeks to reduce data size. A SVR model is then trained at each terminal-node to predict the data locally, and it then trains k non-linear SVR models in a straightforward manner in parallel on multi-core systems. When compared to the normal SVR in LibSVM, our proposed tSVR is more effective in terms of training time and prediction accuracy, according to the numerical test results on datasets from the UCI repository.

3 The Proposed Methodology

In this work, we use regression models, that are presented in [16, 17], for blog feedback prediction. We employ five regression models: K-Nearest Neighbors, Decision Tree Regressor, MLP Regressor, SVR Regressor, and Random Forest Regressor. Figure 1 shows a summary of the regression model for predicting blog feedback. It is divided into three levels: Data pre-prospecting and feature extraction are included in the first level; five different regression models

are contained in the second level's second layer and are used to predict blog feedback along with their underlying principles; and finally, the prediction process is completed in the third level, which includes training, testing, and evaluating the models.

3.1 Data Pre-processing and Feature extraction

The data were preprocessed to remove the missing and unwanted data to obtain a cleaned dataset. Then, the features are extracted from the cleaned data to obtain the best features for training models. We must address two challenges before we can use machine learning for the feedback prediction problem: We require certain data for which the value of the target is already known, and we have to vectorize the instances (blog posts) (train data).

We extract the following features from each page in order to solve the first problem, which is to convert them into vectors:

1. Basic features: How many links and feedbacks increased or decreased in the past (the past is seen relative to baseTime); how many links and feedbacks occurred in the first 24 hours following the publication of the document but before baseTime; how many links and feedbacks occurred in the time period from 48 hours to 24 hours prior to baseTime; combining the aforementioned characteristics by source,
2. Textual features: The most discriminative bag of words features,
3. Weekday features, which include the day of the week the document's primary text was released and the day of the week for which the prediction must be calculated, are binary indicator features,
4. Parent features: We consider a document d_P to be a parent of document d if d is a reply to d_P , that is, if d is referenced by a traceback link on d_P . Parent features include the number of parents, as well as the minimum, maximum, and average number of feedback that the parents have received.

The first problem is resolved by simulating that the chosen date and time would be the current date and time at some point in the past. The chosen time and date are referred to as baseTime.

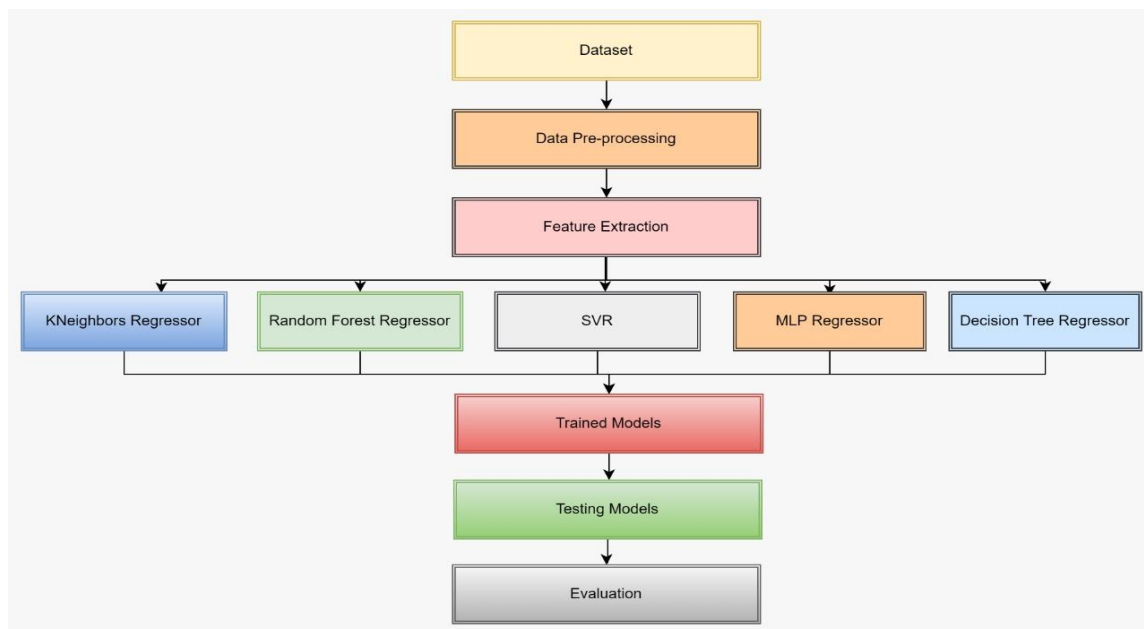


Figure 1: Regression model for the blog feedback prediction

3.2 The proposed models

We proposed a decision tree regressor, MLP regressor, SVR, random forest regressor, and KNeighbors regressor to build five models as presented in [16, 17]. We give a brief overview of these methods as follows.

Decision Trees. Decision trees are used to categorize or predict data using a tree-like model of decisions and their outcomes. For both a classification tree and a prediction, using a decision tree is helpful (regression tree). Given that a single regression tree's dependability is lower than that of a random forest regression, it might not be able to anticipate complex situations (RFR). RFR is the most well-known decision tree-based model and consists of an ensemble of regression trees. [18, 19, 20].

Multi-Layer Perceptron Model. The Multilayer Perceptron (MLP) is a feed forward neural network type that is used for air pollution prediction because it has the capacity to build incredibly complex nonlinear models. The network receives the necessary input parameters, causing the MLP to fire. The input signals produced by these input parameters are transmitted across the network, first from the input layer to the hidden layer and then from the hidden layer to the output layer. The weights, which are real numerical values, are multiplied by the scaled input vector that is introduced by the neurons in the input layer [21, 22, 23].

Support Vector Regression. One use for the 1954-invented support vector machines (SVM) is support vector regression (SVR). The foundations of SVM are structural minimization and statistical learning theory. By transferring the data from a low-dimensional to a high-dimensional space, the kernel function's main goal is to convert the data from a nonlinear to a linear feature space. Radial basis functions, linear basis functions, and polynomial basis functions are examples of common kernel functions that can be utilized in SVR [24, 25, 26].

Random Forest. The ensemble learning-based random forest (RF) machine learning technique can be used to address classification and regression problems. It is a member of the group of algorithms for supervised machine learning. To create the random forest design, a series of trees are simultaneously built using the fundamental approach employed in random forests, known as bagging. Although they are generated, the trees in the random forest do not interact. This method's training phase comprises the creation of several decision trees that may be applied to either classification or regression tasks. The result of random forest in regression problems is the average forecast each tree provides. Because of this, the forecast made using random forest is seen as a mixture of the several predictions made using built-in decision trees [27, 28].

k Nearest Neighbors. One of the most popular machine learning methods is the k Nearest Neighbors (kNN) algorithm. It can be used to categorize. Machine learning uses statistical and mathematical techniques to draw conclusions from the available data. An algorithm for non-parametric learning is kNN. The k nearest neighbors approach organizes new cases in accordance with the requirements of the distance function while saving all of the existing instances. The knowledge used in training is memorized rather than taught. The algorithm's execution results in the determination of a k value. K different factors need to be considered. The distance to the closest k element is determined when a value is reached. The instance is merely assigned to its closest neighbor if k is 1. For the bulk of data sets, the optimal k numbers traditionally varied from 3 to 10. The Euclidean function is frequently used to calculate distances. The Manhattan, Minkowski, and Hamming functions are alternatives to the Euclidean function [29].

The Proposed Ensemble Model. Each regression model in the suggested ensemble has its parameters optimized before being combined with other models to form a single ensemble. This ensemble includes the Decision Tree Regressor, MLP Regressor, SVR, Random Forest Regressor, and KNeighbors Regressor. The best-predicted value that most closely resembles the air quality input attributes is the ensemble model's output.

3.3 Regression Model Training

After the feature extraction step, the features represent the data training for models. Five separate models were subsequently created for testing and training. In order to deploy the models, you must first confirm the prediction value range. If necessary, you should then forecast if the air quality levels are satisfactory or good; if not, the models

and datasets need to be improved once more.

3.4 Evaluation

The most popular evaluation criteria are the correlation coefficient (R2), Mean Absolute Error (MAE), Root Mean Square Error (RMSE), MBE, RRMSE, NSE, and WI. The R2 value shows the fitting degree of regression. MAE represents the difference between predicted and actual values. RMSE focuses on the impact of extreme values based on MAE, while RAE calculates the variance of a model when comparing the performance of different models. As MAE and RMSE depend on the scale of the data that's why RAE can be extremely helpful when comparing different data with different scales.

4 Experimental Results

4.1 Dataset

We use the BlogFeedback dataset. This data comes from blog entries. The blog entries' unprocessed raw HTML files were crawled. The data's prediction objective is to forecast how many comments there will be during the next 24 hours. We pick a basetime (in the past) and select the blog posts that were posted no more than 72 hours prior to the base date/time that we have chosen. After that, using the data that was available at the basetime, we calculate all the features of the chosen blog posts, so each instance is equivalent to a blog post. The target is the amount of comments the blog post receives over the course of the following 24 hours in relation to the basetime.

The basetimes in the train data were between 2010 and 2011. The basetimes of the test data were February and March 2012. This models the situation that might exist in the real world when training data from the past could be used to forecast future events.

The basetimes used to build the train data may have overlapped in time. Therefore, the underlying time intervals might overlap if you merely divide the train into separate pieces. Therefore, to guarantee that the evaluation is fair, you should use the temporally separate train and test splits that are given.

4.2 Results and discussion

We evaluate the proposed models based on the communities and crime dataset. The metrics, that are presented in subsection 3.4, are used to evaluate the proposed model. The results of the proposed model are been shown in Table 1.

Table 1: The results of the regression model

Model	RMSE	MAE	MBE	R	R ²	RRMSE	NSE	WI
Decision Tree	0.0704	0.0524	0.0018	0.94	0.88	24.16	0.88	0.85
Multi-layer Perceptron	0.0495	0.0384	0.0041	0.97	0.94	16.99	0.94	0.89
support vector regression	0.0424	0.0326	0.0048	0.98	0.96	14.55	0.96	0.91
Random forest	0.1040	0.0806	-0.0155	0.87	0.75	35.69	0.74	0.77
K-nearest neighbors	0.0351	0.0242	-0.0014	0.99	0.97	12.06	0.97	0.93

As shown in Table 1, we note that the experimental results enhance the results. It is further observed all the five models produce good results.

In regression results of decision tree, the RMSE, MAE, MBE, R, R2, RRMSE, NSE, and WI terms are reported by 0.0704, 0.0524, 0.0018, 0.94, 0.882, 24.16, 0.88, and 0.85, respectively.

In regression results of multi-layer perceptron, the RMSE, MAE, MBE, R, R2, RRMSE, NSE, and WI terms are reported by 0.0495, 0.0384, 0.0041, 0.97, 0.94, 16.99, 0.94, and 0.89, respectively.

In regression results of SVR, the RMSE, MAE, MBE, R, R2, RRMSE, NSE, and WI terms are reported by 0.0424, 0.0326, 0.0048, 0.98, 0.96, 14.55, 0.96, and 0.91, respectively.

In regression results of random forest, the RMSE, MAE, MBE, R, R2, RRMSE, NSE, and WI terms are reported by 0.1040, 0.0806, -0.0155, 0.87, 0.75, 35.69, 0.74, and 0.77, respectively.

In regression results of K-nearest neighbors, the RMSE, MAE, MBE, R, R², RRMSE, NSE, and WI terms are reported by 0.0351, 0.0242, -0.0014, 0.99, 0.97, 12.06, 0.97, and 0.93, respectively.

We note that the K-nearest neighbors achieved the best results for RMSE, MAE, R, R², RRMSE, NSE, and WI terms. While the best result for MBE is yielded by random forest.

Two experiments of ensembles are conducted in this work including average ensemble and ensemble using K-Neighbors Regressor. These models are evaluated based on the RMSE, MAE, MBE, R, R², RRMSE, NSE, and WI metrics as shown in Table 2.

Table 2: The results of the ensemble models

Model	RMSE	MAE	MBE	R	R ²	RRMSE	NSE	WI
Average Ensemble	0.0463	0.0329	-0.0012	0.97	0.95	12.06	0.95	0.90
Ensemble using KNN regressor	0.0197	0.0205	-9.4427	0.99	0.98	10.19	0.98	0.94

The average ensemble yielded the results of RMSE, MAE, MBE, R, R², RRMSE, NSE, and WI by 0.0463, 0.0329, -0.0012, 0.97, 0.95, 12.06, 0.95, and 0.90, respectively. While ensemble using K-Neighbors regressor achieved the results of RMSE, MAE, MBE, R, R², RRMSE, NSE, and WI by 0.0197, 0.0205, -9.4427, 0.99, 0.98, 10.19, 0.98, and 0.94.

The above results illustrate that our ensemble model is more accurate than the other five regression models for the prediction. The regression results allow to believe that our proposed ensemble model is efficient for handling these data. In addition, the ensemble using KNN regressor achieved the best results for blog feedback prediction.

5 Conclusions

In the last decade, the importance of social media grew unbelievably. In this paper, we proposed different regression models for blog feedback prediction. We aimed to predict the number of feedbacks that blog documents receive. The proposed models include decision tree regressor, MLP regressor, SVR, random forest regressor, and K-Neighbors regressor. The average ensemble and ensemble using K-Neighbors regressor are used for enhancing the proposed models. The MSE, MAE, MBE, R, R², RRMSE, NSE, and WI metrics are used to evaluate the proposed models. The results show that the regression models perform well, which are the competitive models. Also, the ensemble using KNN regressor achieved the best results for blog feedback prediction.

References

- [1] Krisztian Buza. Feedback prediction for blogs. *Data analysis, machine learning and knowledge discovery*. Springer, Cham, 2014. 145-152.
- [2] Harsurinder Kaur, and Husanbir Singh Pannu. Blog response volume prediction using adaptive neuro fuzzy inference system. *2018 9th international conference on computing, communication and networking technologies (ICCCNT)*. IEEE, 2018.
- [3] Krisztian Buza, and Tomáš Horváth. Factorization Machines for Blog Feedback Prediction. *International Conference on Computer Recognition Systems*. Springer, Cham, 2019.
- [4] M. Sterling, P. Leung, D. Wright, and T. F. Bishop. The use of social media in graduate medical education: a systematic review. *Academic Medicine*, vol. 92, no. 7, pp. 1043–1056, 2017.
- [5] E. Bonson, S. Royo, and M. Ratkai. Citizens' engagement on local governments' facebook sites. an empirical analysis: The impact of different media and content types in western Europe. *Government Information Quarterly*, vol. 32, no. 1, pp. 52–62, 2015.
- [6] D. A. Huffaker and S. L. Calvert. Gender, identity, and language use in teenage blogs. *Journal of computer-mediated communication*, vol. 10, no. 2, p. JCMC10211, 2005.
- [7] C. G. Maurya, S. Gore, and D. S. Rajput. A use of social media for opinion mining: An overview (with the use of hybrid textual and visual sentiment ontology). In *Proceedings of International Conference on Recent Advancement on Computer and Communication*, pp. 315–324, Springer, 2018.

- [8] Krisztian Buza, Alexandros Nanopoulos, and Gábor Nagy. Nearest neighbor regression in the presence of bad hubs. *Knowledge-Based Systems* 86 (2015): 250-260.
- [9] Mandeep Kaur, and Prince Verma. "Comment volume prediction using regression." *International Journal of Computer Applications* 151.1 (2016): 1-9.
- [10] Kamaljit Singh, Ranjeet Kaur Sandhu, and Dinesh Kumar. Comment volume prediction using neural networks and decision trees. *IEEE UKSim-AMSS 17th International Conference on Computer Modelling and Simulation, UKSim2015 (UKSim2015)*. 2015.
- [11] Hieu Tran, and Maxim Shcherbakov. Detection and prediction of users attitude based on real-time and batch sentiment analysis of facebook comments. *International conference on computational social networks*. Springer, Cham, 2016.
- [12] Carlos Castillo, et al. Characterizing the life cycle of online news stories using social media reactions. *Proceedings of the 17th ACM conference on Computer supported cooperative work & social computing*. 2014.
- [13] Jingfei Du, Hua Xu, and Xiaoqiu Huang. Box office prediction based on microblog. *Expert Systems with Applications* 41.4 (2014): 1680-1689.
- [14] Md Taufeeq Uddin. Automated blog feedback prediction with ada-boost classifier. *2015 international conference on informatics, electronics & vision (ICIEV)*. IEEE, 2015.
- [15] Minh-Thu Tran-Nguyen, et al. Decision tree using local support vector regression for large datasets. *Asian Conference on Intelligent Information and Database Systems*. Springer, Cham, 2018.
- [16] Hamzah A. Alsayadi, Abdelaziz A. Abdelhamid, El-Sayed M. El-Kenawy, Abdelhameed Ibrahim, and Marwa M. Eid. Improving the Regression of Air Quality Using Ensemble of Machine Learning Models. *Journal of Artificial Intelligence and Metaheuristics* 1.2 (2022): 1-8.
- [17] Hamzah A. Alsayadi, Nima Khodadadi, and Sunil Kumar. Improving the Regression of Communities and Crime Using Ensemble of Machine Learning Models. *Journal of Artificial Intelligence and Metaheuristics* 1.1 (2022): 27-34.
- [18] Wenjuan Wei, Olivier Ramalho, Laetitia Malingre, Sutharsini Sivanantham, John C Little, and Corinne Mandin. Machine learning and statistical models for predicting indoor air quality. *Indoor Air*, 29(5):704– 726, 2019.
- [19] El-Sayed M. El-Kenawy, Seyedali Mirjalili, Fawaz Alassery, Yu-Dong Zhang, Marwa Metwally Eid, Shady Y. El-Mashad, Bandar Abdullah Aloyaydi, Abdelhameed Ibrahim, and Abdelaziz A. Abdelhamid. Novel Meta-Heuristic Algorithm for Feature Selection, Unconstrained Functions and Engineering Problems. *IEEE Access*, 10:40536–40555, 2022.
- [20] Abdelaziz A. Abdelhamid, El-Sayed M. El-Kenawy, Bandar Alotaibi, Ghada M. Amer, Mahmoud Y. Abdelkader, Abdelhameed Ibrahim, and Marwa Metwally Eid. Robust Speech Emotion Recognition Using CNN+LSTM Based on Stochastic Fractal Search Optimization Algorithm. *IEEE Access*, 10:49265– 49284, 2022.
- [21] S Abdullah, M Ismail, and AN Ahmed. Multi-layer perceptron model for air quality prediction. *Malaysian Journal of Mathematical Sciences*, 13:85–95, 2019.
- [22] Doaa Sami Khafaga, Amel Ali Alhussan, El-Sayed M. El-kenawy, Ali E. Takieldeem, Tarek M. Hassan, Ehab A. Hegazy, Elsayed Abdel Fattah Eid, Abdelhameed Ibrahim, and Abdelaziz A. Abdelhamid. Metaheuristics for Feature Selection and Classification in Diagnostic Breast-Cancer. *Computers, Materials & Continua*, 73(1):749–765, 2022.
- [23] Doaa Sami Khafaga, Amel Ali Alhussan, El-Sayed M. El-kenawy, Abdelhameed Ibrahim, Said H. Abd Elkhaliq, Shady Y. El-Mashad, and Abdelaziz A. Abdelhamid. Improved Prediction of Metamaterial Antenna Bandwidth Using Adaptive Optimization of LSTM. *Computers, Materials & Continua*, 73(1):865–881, 2022.
- [24] Ruizhi Zhong, Raymond L Johnson Jr, and Zhongwei Chen. Using machine learning methods to identify coals from drilling and logging-while-drilling lwd data. In *Asia Pacific Unconventional Resources Technology Conference, Brisbane, Australia, 18-19 November 2019*, pages 970–994. Unconventional Resources Technology Conference, 2020.
- [25] Nagwan Abdel Samee, El-Sayed M. El-Kenawy, Ghada Atteia, Mona M. Jamjoom, Abdelhameed Ibrahim, Abdelaziz A. Abdelhamid, Noha E. El-Attar, Tarek Gaber, Adam Slowik, and Mahmoud Y. Shams. Metaheuristic Optimization Through Deep Learning Classification of COVID-19 in Chest X-Ray Images. *Computers, Materials & Continua*, 73(2):4193–4210, 2022.
- [26] Hussah Nasser AlEisa, El-Sayed M. El-kenawy, Amel Ali Alhussan, Mohamed Saber, Abdelaziz A. Abdelhamid, and Doaa Sami Khafaga. Transfer Learning for Chest X-rays Diagnosis Using Dipper Throated Algorithm. *Computers, Materials & Continua*, 73(2):2371–2387, 2022.
- [27] Zhi-Hua Zhou. *Machine learning*. Springer Nature, 2021.

- [28] Doaa Sami Khafaga, Amel Ali Alhussan, El-Sayed M. El-Kenawy, Abdelhameed Ibrahim, Marwa Metwally Eid, and Abdelaziz A. Abdelhamid. Solving optimization problems of metamaterial and double t-shape antennas using advanced meta-heuristics algorithms. *IEEE Access*, 10:74449–74471, 2022.
- [29] Burhan BARAN. Air quality index prediction in besiktas district by artificial neural networks and k nearest neighbors. *Muhendislik Bilimleri ve Tasarım Dergisi* "", 9(1):52–63, 2021.