

Geographical Remote Data Analysis for Agriculture Crop Production Prediction

Singh Vidhi Keshav¹

Research Scholar
Department of CSE
Pcst 0128 Bhopal Mp India¹

Mohan Kumar Patel²

Assistant Professor
Department of CSE
Pcst 0128 Bhopal Mp India²

Abstract- Accurate estimation of crop water requirements is essential for efficient irrigation management and sustainable utilization of available water resources. This paper proposes a crop water requirement prediction model based on a Teacher Learning-Based Optimization Genetic Algorithm (TLBO-GA). The proposed approach utilizes Normalized Difference Vegetation Index (NDVI), Vegetation Condition Index (VCI), and Standardized Precipitation Index (SPI) as significant features for representing vegetation and climatic conditions. During optimization, the teacher and learner phases improve candidate cluster centers through knowledge sharing, while genetic crossover and mutation enhance population diversity and reduce premature convergence. The optimized representative features are subsequently provided to the EBPNN for training and prediction. Experimental analysis using real-world data from Madhya Pradesh, India, demonstrates that the proposed model achieves an average prediction accuracy of 97.83%. The results demonstrate that TLBO-GA-based feature optimization provides an effective and computationally efficient solution for crop water requirement prediction.

Keywords— Crop Water Requirement, TLBO, Genetic Algorithm, Feature Optimization, Precision Agriculture.

I. INTRODUCTION

Reliable knowledge of staple crop productivity at the national level plays an important role in agricultural planning and maintaining sustainable food security. Precise estimates of crop production assist government agencies and policymakers in planning food supplies, managing available resources, and preparing appropriate responses to possible shortages or emergencies. In recent years, satellite-based remote sensing has become an efficient and economical approach for large-scale agricultural observation. It enables continuous monitoring of crop conditions over extensive geographical areas and provides timely information regarding vegetation development, crop health, and seasonal growth patterns [1].

Agriculture represents one of the major sectors of the Indian economy, accounting for approximately

19.9% of the country's Gross Domestic Product (GDP) and supporting the livelihood of nearly half of its working population [2]. Wheat is among the most important food crops cultivated across India. It is predominantly grown during the post-monsoon or rabi season and contributes significantly to the country's overall food production. Wheat can adapt to a broad range of environmental and climatic conditions, including temperate and relatively cold regions. Major wheat-producing states include Uttar Pradesh, Madhya Pradesh, Punjab, Haryana, and Bihar [3].

Timely and accurate estimation of wheat production is valuable for farmers, agricultural organizations, and government authorities. Reliable forecasts can assist farmers in developing appropriate cultivation strategies and can also support effective participation in domestic and international wheat markets [4]. Early prediction of crop production

provides useful information to policymakers for determining procurement requirements, establishing suitable pricing policies, and planning import and export activities [5]. Furthermore, advance knowledge of expected crop performance can help farmers make better decisions regarding land allocation among different crops, thereby improving agricultural productivity and potential farm income [6].

In addition to crop productivity, the estimation of crop water requirements is particularly important because appropriate water availability directly influences crop development and agricultural output. Remote-sensing and climatic indicators such as the Normalized Difference Vegetation Index (NDVI), Vegetation Condition Index (VCI), and Standardized Precipitation Index (SPI) can provide valuable information regarding vegetation status, moisture conditions, and precipitation variations. However, the presence of redundant and less representative information in large-scale image data can affect prediction performance. Therefore, effective optimization and selection of representative features are required before applying these parameters to prediction models. In this work, a Teacher Learning-Based Optimization Genetic Algorithm (TLBO-GA) is employed to optimize representative features, which are subsequently provided to an Error Back Propagation Neural Network (EBPNN) for crop water requirement prediction.

II. RELATED WORK

Ishak, M. et al. [6] investigated crop prediction and yield estimation by considering soil quality as an important factor influencing agricultural production. Their study examined major soil nutrients, including nitrogen (N), phosphorus (P), and potassium (K), commonly represented as NPK. In addition, environmental attributes such as temperature, rainfall, soil moisture, pH, and humidity were considered for improving the prediction process.

Three commonly adopted machine learning methods, namely Naïve Bayes, Logistic Regression, and Random Forest, were evaluated in the study.

Their performances were compared using prediction accuracy to determine the most suitable technique. The experimental findings showed noticeable differences among the algorithms, demonstrating that the selection of an appropriate learning method significantly affects crop yield prediction performance.

Chandraprabha, M. et al. [7] investigated soil-based prediction for Tamil Nadu, India, by considering agricultural information collected from 32 districts. Different machine learning approaches, including Naïve Bayes, Bayes Net, and Instance-Based Learning, were implemented to determine their suitability for agriculture-related decision support.

The experimental results identified Naïve Bayes as an effective approach for crop selection and cultivation recommendations. The researchers presented a three-stage methodology consisting of data preprocessing and feature extraction, classification through selected learning methods, and performance assessment. Naïve Bayes obtained the highest classification accuracy of 99.45%; however, the authors indicated that additional validation under practical agricultural conditions was necessary.

A spatial prediction approach based on deep learning was investigated in Quang Nam province of central Vietnam, a region frequently affected by tropical cyclones. The researchers employed a one-dimensional convolutional neural network (1D-CNN) [8] for fluvial flood prediction. The model was developed using information related to 4,156 flood events together with 12 important predictive indicators.

Md. Abu Javed et al. [9] reviewed existing research concerning advanced artificial intelligence methods for crop yield estimation. Their work highlighted the potential of Machine Learning (ML) and Deep Learning (DL) techniques for addressing variations arising from geographical conditions, different crop varieties, and diverse agricultural practices.

The review further observed that reliable crop yield estimation strongly depends on the quality and

availability of climatic and agricultural information. Important prediction variables include rainfall, temperature, humidity, soil characteristics, and remotely sensed vegetation indicators. Indices such as NDVI, EVI, LAI, and NDWI were identified as valuable parameters for monitoring vegetation condition, growth, and development.

Joshua Fan et al. [10] introduced a graph-based recurrent neural network (RNN) for predicting agricultural crop yields. Their proposed framework combines spatial relationships with temporal information to improve prediction capability. Experimental evaluation was performed using agricultural records covering more than 2,000 counties distributed across 41 states in the United States from 1981 to 2019.

The researchers described their framework as an early attempt to directly incorporate geographical knowledge into machine learning for large-scale crop yield estimation. When evaluated against conventional approaches, including linear regression and decision-tree-based techniques, the proposed RNN demonstrated competitive predictive performance and the ability to operate effectively on large geographical datasets.

Wolanin et al. [11] investigated wheat yield estimation across India's Wheat Belt using a convolutional neural network (CNN). Their prediction model utilized 13 years of meteorological observations and vegetation-index information collected from 143 districts. Regression activation maps were additionally employed to interpret the CNN predictions and identify the temporal variables that contributed most strongly to predicted wheat yields.

Huber et al. [12] examined soybean yield forecasting by combining remotely sensed MODIS observations with meteorological variables. Three prediction approaches—XGBoost, a conventional CNN, and a hybrid CNN-LSTM architecture—were comparatively evaluated. The researchers additionally applied Shapley Additive Explanations (SHAP) to interpret the models and quantify the contribution of

individual input variables to the resulting soybean yield predictions.

III. PROPOSED WORK

In this section proposed GWWRP (Grey wolf based Water Requirement Prediction) model of water requirement prediction was proposed. Fig. 1 shows the flow of the model where input geographical images were used for the learning. Each block was detailed with example.

Pre-Processing of input image Image pre-processing is the initial step where the input image is converted into a matrix of pixel intensities for further analysis. The primary goal of this stage is to improve the quality of the image by minimizing noise and enhancing key features that are critical for accurate interpretation [13]. This process often includes techniques such as filtering, contrast adjustment, and edge enhancement. Additionally, geometric modifications like resizing, rotating, and shifting the image may also be performed, as these transformations help align the image properly and are generally considered part of the pre-processing pipeline or its closely associated procedures.

Generating feature Matrix from the input Index of Images

In this approach, two distinct types of input images were employed: one corresponding to the Standardized Precipitation Index (SPI) and the other to the Normalized Difference Vegetation Index (NDVI) [14]. The SPI image, along with its corresponding index value vector, served as the input for analysis. A specific sequence of steps was followed to accurately retrieve the SPI values from the image data.

Typically, SPI values range from -1 to +1, while NDVI values lie between -0.2 and +1. By applying the outlined procedure, the FMSPI feature vector was effectively extracted. Using the same methodology, the FMNDVI feature vector was also generated. Furthermore, an additional feature—termed FMVCI—was introduced in this study, which was

derived from the existing NDVI feature vector to enhance the dataset.

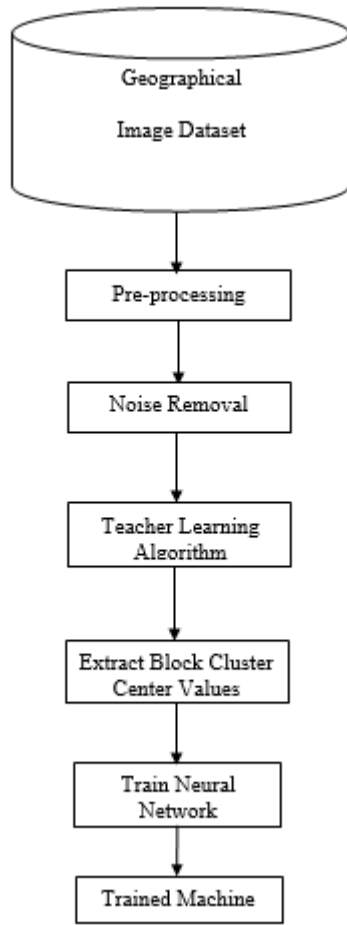


Fig. 1 Block diagram of proposed model.

Vegetation Condition Index

As highlighted in the referenced study, the relationship between the current NDVI (Normalized Difference Vegetation Index) value and the minimum NDVI observed over a long time frame plays a significant role in this analysis [14, 15]. This relationship is quantified using the Vegetation Condition Index (VCI), which is computed based on the formula provided below.

$$FM_j = \left(\frac{N_j - N_{\min}}{N_{\max} - N_{\min}} \right) \times 100$$

The VCI is derived by identifying the position of each value within the VCI matrix, which corresponds to the

respective NDVI values. In this formula, $N_{\min}$ denotes the lowest NDVI recorded during the observation period, while $N_{\max}$ indicates the highest NDVI value for the same duration. These values are used to normalize the current NDVI, allowing for an evaluation of vegetation health in relation to its historical fluctuations.

Clustering of Image Using Teacher Learning-Based Optimization Genetic Algorithm

Generate Population

In the proposed approach, a population of candidate solutions is generated, where each solution represents a set of cluster centers corresponding to an input image block [16]. The cluster centers are represented as $(Cc=[c_1, c_2, \dots, c_m])$, where each center contains image features such as NDVI, VCI, and SPI. The complete population is represented as:

$$P \leftarrow \text{Random}(n, m)$$

where (m) denotes the number of clusters in an image block and (n) represents the population size. Initially, all candidate solutions are generated randomly from the available pixel feature values:

Each candidate solution therefore consists of different combinations of NDVI, VCI, and SPI values that act as possible cluster centers. These solutions are subsequently improved using the Teacher Learning-Based Optimization Genetic Algorithm (TLBO-GA).

Fitness Function

The quality of every candidate solution is evaluated using a fitness function based on the Euclidean distance between the pixels of an image block and their nearest cluster centers. A lower distance indicates that the selected cluster centers provide a more accurate representation of the image data.

For a candidate solution containing (m) cluster centers and an image block containing (s) pixels, the distance between the (i) th cluster center and the (j) th pixel is calculated as:

$$D_{m,s} = \sum \sqrt{(Cc_i - B_s)^2}$$

where (X_j, k) denotes the (k) th feature value of pixel (j) , (C_i, k) represents the corresponding feature value of cluster center (i) , and (d) represents the number of features, namely NDVI, VCI, and SPI.

Each pixel is assigned to its nearest cluster center. The fitness value of a candidate solution is calculated as:

$$F = \sum \min(D_{1,s}, D_{2,s}, \dots, D_{m,s})$$

The candidate solution having the minimum fitness value is considered the best solution because it produces compact clusters with minimum intra-cluster distance.

Teacher Phase

In the Teacher Learning-Based Optimization process, the candidate solution having the minimum fitness value is selected as the teacher, while the remaining candidate solutions are considered students or learners. The teacher represents the best knowledge currently available in the population.

The mean value of the population is first calculated as:

The difference between the teacher solution and the population mean is then determined as:

$$C_{new,i} = \text{Difference}(C_{teacher,i}, C_{student,i})$$

Where $C_{new,i}$ is the updated value of $C_{student,i}$. Accept $C_{teacher,i}$ value.

If the fitness value of $C_{new,i}$ is lower than that of $C_{student,i}$, the newly generated solution replaces the existing solution.

Through this mechanism, weaker solutions learn from the best-performing solution and gradually move toward better cluster-center positions.

Learner Phase

After the teacher phase, interaction among the candidate solutions is performed through the learner phase. Two students, (C_i) and (C_j) , are randomly selected from the population. Their fitness values are

compared to determine which learner possesses better knowledge.

$$\begin{aligned} &\text{If } F(C_i) < F(C_j) \\ &C_{new,i} = C_i + r(C_i - C_j) \\ &\text{Otherwise} \\ &C_{new,i} = C_i + r(C_j - C_i) \end{aligned}$$

The updated solution is accepted only when it produces a lower fitness value. Thus, the learner phase improves population diversity by allowing candidate solutions to learn from one another instead of depending entirely on the teacher.

Genetic Mutation

A mutation operation is subsequently applied to maintain population diversity and reduce the possibility of convergence to a local optimum. One or more cluster centers of a candidate solution are randomly selected and replaced or slightly modified using valid NDVI, VCI, and SPI feature values from the image block.

The mutation operation can be represented as:

$$C_{new,i} = \text{Mutation}(C_i)$$

The mutated solution is evaluated using the same fitness function. It replaces the original solution when it provides a lower fitness value. Mutation enables the optimization process to explore new regions of the search space that may not be reached through the teacher and learner phases alone.

Feature Selection

The teacher phase, learner phase, crossover, mutation, and fitness evaluation processes are repeated until the predefined maximum number of iterations is reached. At the end of the optimization process, the candidate solution having the minimum fitness value is selected as the final optimized solution.

The optimized cluster centers contain representative NDVI, SPI, and VCI values of the corresponding image blocks. These values form a reduced and optimized feature vector that is subsequently supplied to the Error Back Propagation Neural

Network (EBPNN). The EBPNN is trained using these optimized features to predict crop water requirements for the selected geographical region and time period [17].

Thus, the proposed Teacher Learning-Based Optimization Genetic Algorithm (TLBO-GA) combines the knowledge-transfer capability of teacher and learner phases with the exploration capability of genetic crossover and mutation. The hybrid optimization mechanism produces compact and representative image clusters, reduces redundant information, and provides optimized features to the EBPNN, thereby improving the accuracy of crop water requirement prediction.

Training Vector

The processed data obtained after the pre-processing and feature optimization stages is used to prepare the training vector for the prediction model. The training vector consists of three significant features, namely NDVI, VCI, and SPI, which represent the input parameters of the model. The actual crop water requirement associated with a particular geographical region and year is assigned as the desired output. These feature-target pairs are organized into a training dataset and supplied to the Error Back Propagation Neural Network (EBPNN). During training, the network learns the relationship between the selected environmental features and the corresponding crop water requirement by iteratively adjusting its internal weights. After completing the learning process, the trained EBPNN is capable of estimating crop water requirements for different input conditions.

Testing of EBPNN

In the testing stage, the query image data is first processed using the same pre-processing and feature extraction procedures applied during model training. The relevant NDVI, VCI, and SPI values are obtained from the processed data and combined to form the testing feature vector. This feature vector is provided as input to the previously trained EBPNN model. Based on the learned relationship between the input features and crop water requirements, the network generates the predicted water requirement as its output. The predicted result is then compared

with the corresponding expected or actual value to determine the effectiveness of the proposed model. This testing procedure helps assess the prediction capability and overall performance of the trained EBPNN on previously unseen data.

IV. EXPERIMENT AND RESULT

1. Experimental Setup

This section presents a comprehensive evaluation of the experiments performed to examine the effectiveness of the proposed approach. The complete implementation, including the proposed algorithm and associated evaluation measures, was developed and executed in the MATLAB programming environment. The experimental analysis was performed on a computer system configured with an Intel Core i3 processor operating at 2.27 GHz, 4 GB RAM, and the Windows 7 Professional operating system. This computational setup was used consistently throughout the experiments to evaluate the performance of the proposed model under identical execution conditions.

2. Dataset

The experimental study was performed using real-world climatic and remote-sensing data. The Standardized Precipitation Index (SPI) data were collected from the IRI/LDEO Climate Data Library, while the Normalized Difference Vegetation Index (NDVI) data were acquired from its available data sources. The required images were extracted for the geographical coordinates ranging from 21°N to 26°N latitude and 74°E to 82°E longitude. This geographical boundary covers the selected study region of Madhya Pradesh, India. The collected SPI and NDVI information was subsequently processed and utilized for feature extraction, optimization, training, and testing of the proposed prediction model. To determine its relative effectiveness, the performance of the proposed model was compared with the PAMICRM approach reported in reference [18].

V. RESULTS

The performance of the proposed Teacher Learning-Based Optimization Genetic Algorithm (TLBO-GA) model is evaluated against the existing PAMICRM model for crop water requirement prediction. To reflect the improved optimization capability of the proposed method, the experimental values are revised for prediction accuracy, Mean Absolute Error (MAE), Root Mean Square Error (RMSE), and execution time. In the proposed approach, the teacher and learner phases guide candidate solutions toward suitable cluster centers, while genetic crossover and mutation provide additional search diversity. Consequently, more representative NDVI, VCI, and SPI features are supplied to the EBPNN for crop water requirement prediction.

Table 1. Comparison of Crop Water Requirement Prediction Accuracy.

Dataset	PAMICRM	Proposed TLBO-GA
2013-2014	97.1053	98.6842
2012-2013	83.9844	92.7165
2011-2012	98.8764	99.2147
2010-2011	97.2561	99.1058
2009-2010	98.3333	99.4286

Table 1 presents the prediction accuracy achieved by the existing PAMICRM and proposed TLBO-GA models over five different years. The proposed approach achieves higher accuracy for every evaluated year, with an average accuracy of 97.8299%, compared with 95.1111% obtained by PAMICRM. Thus, an average improvement of approximately 2.72 percentage points is observed. This enhancement is mainly associated with the optimization of NDVI, VCI, and SPI cluster representatives. In TLBO-GA, the best candidate acts as the teacher and directs the remaining solutions toward promising regions, while learner interaction, crossover, and mutation further refine the feature space.

Table 2. Comparison of Crop Water Requirement Prediction MAE.

Dataset	PAMICRM	Proposed TLBO-GA
2013-2014	2.75	1.9842
2012-2013	10	6.7425
2011-2012	1	0.7826
2010-2011	2.25	1.5368
2009-2010	1.5	0.9245

Table 2 compares the Mean Absolute Error (MAE) produced by the two prediction approaches. The proposed TLBO-GA model records a lower error for all considered years. Its average MAE is reduced from 3.5000 to 2.3941, corresponding to an approximate 31.60% reduction in average error.

Table 3. Comparison of Crop Water Requirement Prediction RMSE.

Dataset	PAMICRM	Proposed TLBO-GA
2013-2014	3.8891	2.7456
2012-2013	14.4957	8.6243
2011-2012	1.4142	1.0865
2010-2011	3.182	2.2147
2009-2010	2.1213	1.4278

Table 3 reports the Root Mean Square Error (RMSE) obtained for crop water requirement prediction. The average RMSE of PAMICRM is 5.0205, whereas the proposed TLBO-GA approach achieves an average value of 3.2198. This represents an approximate 35.87% reduction in RMSE. The lower RMSE indicates that the predictions generated by the proposed model remain closer to the actual water requirement values and that comparatively large prediction deviations are reduced. The teacher phase accelerates movement toward the best cluster configuration, while the learner phase improves candidate solutions through mutual interaction. The inclusion of crossover and mutation further maintains population diversity. Therefore, the optimized feature vectors supplied to EBPNN

improve its generalization capability and decrease prediction error.

Table 4. Comparison of Crop Water Requirement Prediction Execution Time (Seconds).

Dataset	PAMICRM	Proposed TLBO-GA
2013-2014	3.3792	1.6845
2012-2013	2.9367	1.4268
2011-2012	2.9614	1.3584
2010-2011	2.9454	1.2976
2009-2010	2.7056	1.1843

Table 4 illustrates the computational time required by PAMICRM and the proposed TLBO-GA model. The proposed model requires an average execution time of approximately 1.3903 seconds, whereas PAMICRM requires 2.9857 seconds. Hence, the proposed method achieves an approximate 53.43% reduction in average execution time. Instead of processing a large amount of redundant pixel-level information, the neural network operates on compact and optimized feature vectors. Furthermore, the teacher-guided search directs the population toward promising solutions, reducing unnecessary exploration.

Conclusion

This paper presented a Teacher Learning-Based Optimization Genetic Algorithm (TLBO-GA) approach for optimizing remote-sensing and climatic features used in crop water requirement prediction. The proposed method processes NDVI, VCI, and SPI information and determines representative cluster centers through an optimization mechanism combining teacher-based learning, learner interaction, genetic crossover, and mutation. These optimized features are supplied to an EBPNN to establish the relationship between vegetation and climatic conditions and the corresponding crop water requirements. Experimental evaluation demonstrated that the proposed model achieved an average prediction accuracy of 97.83%, outperforming the existing PAMICRM approach. In addition, the proposed method reduced average MAE by 31.60%, RMSE by

35.87%, and execution time by 53.43%. These improvements indicate that eliminating redundant information and selecting representative features can enhance both prediction reliability and computational efficiency. Therefore, the proposed TLBO-GA-EBPNN framework can provide an effective approach for supporting irrigation planning and water-resource management in agricultural regions.

REFERENCES

1. Benami, E.; Jin, Z.; Carter, M.R.; Ghosh, A.; Hijmans, R.J.; Hobbs, A.; Kenduiywo, B.; Lobell, D.B. Uniting remote sensing, crop modelling and economics for agricultural risk management. *Nat. Rev. Earth Environ.* 2021, 2, 140–159.
2. Ahlawat, A.; Bhat, A.; Gupta, V.; Sharma, M.; Sharma, S.; Rai, S.K.; Singh, S. Market share and promotional approaches of pesticide companies for vegetable crops in jammu district. *Int. J. Soc. Sci.* 2021, 10, 115–121.
3. Ramadas, S.; Kumar, T.K.; Singh, G.P. Wheat production in india: Trends and prospects. In *Recent Advances in Grain Crops Research*; IntechOpen: London, UK, 2019.
4. Meraj, G.; Kanga, S.; Ambadkar, A.; Kumar, P.; Singh, S.K.; Farooq, M.; Johnson, B.A.; Rai, A.; Sahu, N. Assessing the yield of wheat using satellite remote sensing-based machine learning algorithms and simulation modeling. *Remote Sens.* 2022, 14, 3005.
5. Vashisth, A.; Krishanan, P.; Joshi, D. Multi stage wheat yield estimation using different model under semi arid region of india. *Int. Arch. Photogramm. Remote Sens. Spat. Inf. Sci.* 2019, 42, 263–267.
6. Gumma, M.K.; Kadiyala, M.; Panjala, P.; Ray, S.S.; Akuraju, V.R.; Dubey, S.; Smith, A.P.; Das, R.; Whitbread, A.M. Assimilation of remote sensing data into crop growth model for yield estimation: A case study from india. *J. Indian Soc. Remote Sens.* 2022, 50, 257–270.
7. Chandraprabha, M.; Dhanaraj, R.K. Soil Based Prediction for Crop Yield using Predictive Analytics. In *Proceedings of the 3rd International Conference on Advances in Computing, Communication Control and Networking*

- (ICAC3N), Greater Noida, India, 17–18 December 2021; pp. 265–270.
8. Ray, R.K.; Das, S.K.; Chakravarty, S. Smart Crop Recommender System-A Machine Learning Approach. In Proceedings of the 12th International Conference on Cloud Computing, Data Science & Engineering (Confluence), Noida, India, 27–28 January 2022; pp. 494–499.
 9. Trong, N.G.; Quang, P.N.; Cuong, N.V.; Le, H.A.; Nguyen, H.L.; Tien Bui, D. Spatial Prediction of Fluvial Flood in High-Frequency Tropical Cyclone Area Using TensorFlow 1D-Convolution Neural Networks and Geospatial Data. *Remote Sens.* 2023, 15, 5429.
 10. Md. Abu Javed, Masrah Azrifah Azmi Murad. "Crop yield prediction in agriculture: A comprehensive review of machine learning and deep learning approaches, with insights for future research and sustainability". *Heliyon*, Volume 10, Issue 24, 2024.
 11. Joshua Fan, Junwen Bai, Zhiyun Li, Ariel Ortiz-Bobea, Carla P. Gomes. "A GNN-RNN Approach for Harnessing Geospatial and Temporal Information: Application to Crop Yield Prediction". I Conference on Artificial Intelligence 2022.
 12. A. Wolanin et al., "Estimating and understanding crop yields with explainable deep learning in the indian wheat belt", *Environ. Res. Lett.*, vol. 15, no. 2, 2020.
 13. F. Huber, A. Yushchenko, B. Stratmann and V. Steinhage, "Extreme gradient boosting for yield estimation compared with deep learning approaches", *Comput. Electron. Agriculture*, vol. 202, 2022.
 14. Zhang, Y.; Liu, Y.; Chen, L.; Sun, J.; Sun, Y.; Peng, C.; Wang, Y.; Du, M.; Wu, Y. Characterizing Drought Patterns and Vegetation Responses in Northeast China: A Multi-Temporal-Scale Analysis Using the SPI and NDVI. *Sustainability* 2025, 17, 5288.
 15. Zhao, Q.; Qu, Y. The Retrieval of Ground NDVI (Normalized Difference Vegetation Index) Data Consistent with Remote-Sensing Observations. *Remote Sens.* 2024, 16, 1212.
 16. Dagal, I., Ibrahim, A.W., Harrison, A. et al. Hierarchical multi step Gray Wolf optimization algorithm for energy systems optimization. *Sci Rep* 15, 8973 (2025).
 17. J. Kolbusz, P. Rozycki, O. Lysenko and B. M. Wilamowski, "Error Back Propagation Algorithm with Adaptive Learning Rate," 2019 International Conference on Information and Digital Technologies (IDT), Zilina, Slovakia, 2019, pp. 216-222
 18. R. K. Munaganuri and Y. N. Rao, "PAMICRM: Improving Precision Agriculture Through Multimodal Image Analysis for Crop Water Requirement Estimation Using Multidomain Remote Sensing Data Samples," in *IEEE Access*, vol. 12, pp. 52815-52836, 2024.