

Prediction of Grain Yield in Anhui Province Based on Multi-Source Remote Sensing Data Introduction

Mengxin Ji, Haifeng Jiang*

School of Business, Anhui University of Technology, Ma'anshan 243032, Anhui, China

**Author to whom correspondence should be addressed.*

Copyright: © 2026 Author(s). This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY 4.0), permitting distribution and reproduction in any medium, provided the original work is cited.

Abstract: Grain output forecasting is an important means to ensure national food security, and accurate output forecasts can provide scientific support for government decision-making. This paper takes Anhui Province as the research area, builds a Stacking ensemble learning model based on multi-source remote sensing data, and predicts grain output in Anhui. The study uses annual structured features (arable land area, sown area, electricity consumption, fertilizer usage, etc.) and monthly time series features (NDVI integral, NDVI mean, sunshine hours, etc.) to build a Stacking ensemble learning prediction model, integrating XGBoost, MLP, and RandomForest base models. XGBoost uses annual features to capture structured information at the annual scale, MLP uses monthly features to capture temporal variation patterns within the growing season, and RandomForest uses fusion features to capture interactions between features. By adaptively learning the weights of each base model through the Ridge regression meta-learner, effective integration of multi-source information is achieved. Experimental results show that the stacking ensemble model performs well under the retain-one cross-validation method: the determination coefficient reaches 0.9575, RMSE is 140.90 kg/ha, and average absolute percentage error MAPE is 4.08%. Compared to the XGBoost benchmark model that uses only annual features, the stacking model improves by 0.0039, and RMSE decreases by 6.21 kg/ha, validating the advantages of multi-source data fusion. Feature importance analysis showed that sown area and arable land area were the most critical factors affecting yield forecasting, together accounting for 60.4% of the characteristic importance. Meta-learner weight analysis shows that XGBoost is the main contributor (weight 1.0224), while MLP and RandomForest serve as auxiliary contributors (weights 0.0309 and 0.0230, respectively), providing dynamic growth season information that annual characteristics cannot capture. The research results of this paper provide a new method for grain output forecasting in Anhui Province, validating the application value of multi-source remote sensing data fusion in yield forecasting, which can offer scientific support for government decision-making and promote sustainable agricultural development.

Keywords: Grain production forecasting; Multi-source remote sensing data; Stacking integrated learning; XGBoost; MLP

Online publication: Sept 3, 2026

1. Introduction

Grain has always been the foundation of national survival and civilizational progress, and it is also a critical resource for maintaining social stability and driving economic growth. The stability and security of grain production are directly related to the high-quality development of the national economy and the comprehensive advancement of society. As one of the major grain-producing regions in China, the southern part of the country is characterized by complex and variable climatic conditions, considerable differences in soil types, a high degree of intensive agricultural production, and significant sensitivity to climatic disturbances. In this context, Anhui Province, as a major agricultural province, plays an important role in safeguarding national food security, regulating regional supply and demand, and stabilizing the national grain market. Fluctuations in grain yield therefore have profound implications for the national food security landscape and the sustainable development of agriculture.

Domestic and international scholars have conducted extensive research in the field of grain production and forecasting^[1], and most studies have found that, compared with single models^[2], ensemble models can improve predictive performance by leveraging the complementary advantages of different models^[3]. Jia Mengqi et al. combined ARIMA, GRNN, and LSTM models to forecast grain yield in Baoding City, Hebei Province, and found that the combination of ARIMA and GRNN achieved the best predictive performance^[4]. Shen Y et al. investigated the integration of LSTM with the random forest (RF) model and demonstrated that this hybrid model could better capture both time-series characteristics and nonlinear features^[5]. Wang X et al. further constructed a dual-branch deep learning model, using LSTM to process meteorological and remote sensing data on one branch and CNN to simulate soil characteristics on the other, confirming that the model accurately integrates different data sources while maintaining sufficient prediction accuracy^[6]. Zhao Yayue et al. proposed a combined forecasting scheme based on support vector regression (SVR) and a CNN-LSTM-attention mechanism for grain yield in Hebei Province^[7]. Ma Dianjing et al. used the Stacking ensemble algorithm to predict grain yield in southern China, providing an important reference for this study^[8]. Although the above studies have achieved favorable results in specific scenarios, as data become increasingly complex, how to leverage the advantages of multiple models further to optimize prediction remains a challenge that needs to be addressed in future research.

Subsequent scholars incorporated multi-source remote sensing data into their studies, steering grain yield prediction in a new direction^[9]. Basso et al. added remotely sensed spatial measurements to traditional indicators and achieved a prediction error of only 10%, demonstrating the suitability of remote sensing data for grain yield forecasting and providing an effective reference for subsequent researchers^[10]. Since then, the integration of remote sensing data and environmental features has been widely adopted in ensemble models, especially in processing soil and meteorological data^[11]. Panov et al. found that combining Sentinel-1B radar satellite-derived remote sensing parameters with ground observation data could improve the accuracy of cereal yield predictions^[12]. Building on this, Zhang Lixin used a GCN-DSTAMNet ensemble model with multi-source remote sensing data to predict grain yield, achieving promising results and providing important inspiration for the present study^[13].

Employing scientific models for predictive analysis can not only provide data support for policy-making but also promote the transformation of agricultural production toward sustainability and efficiency, thereby contributing to rural economic development and social stability, and fostering synergies among food security, ecological security, and economic growth^[14]. However, with the continued intensification of risks such as climate change, the global food security situation is becoming increasingly complex and urgent, and such

uncertainty poses a serious threat to the agricultural sector^[15]. Under these circumstances, using multi-source remote sensing data to predict future grain yield in Anhui Province can both promptly identify and correct unreasonable issues in local grain production and provide a reference for grain production decisions in other provinces and at the national level, thereby ensuring the safety and stability of the food supply.

Most existing studies have adopted feature modeling based solely on a single temporal scale, without exploiting the complementary information embedded in multi-scale data. To address this gap, this paper proposes a dual-scale feature fusion strategy that combines annual statistical features with monthly remote sensing features: annual data reflect the basic conditions of agricultural production, while monthly data capture the dynamic changes in crop growth, and the two types of data complement each other at the model level. In addition, current studies tend to use either a single machine learning model or a simple weighted averaging ensemble approach, which often fails to realize the advantages of different models fully. To this end, this study constructs a three-layer Stacking ensemble learning framework—the base model layer comprises three heterogeneous models, namely XGBoost, MLP, and RF, each handling different types of features; the meta-learner layer employs Ridge regression to adaptively learn the optimal weight combination, thereby integrating the strengths of all models. Grain yield prediction often faces the challenge of limited sample sizes, and most existing studies lack targeted feature engineering and evaluation strategies, which can easily lead to the curse of dimensionality and data leakage. In this paper, we extract growing-season features with clear agronomic significance to reduce dimensionality, adopt a combination of leave-one-out cross-validation and time-series split to prevent leakage of future information, and then screen the core features through analysis of variance and correlation analysis. This approach enhances the model's generalization ability and computational efficiency, enabling a more accurate analysis of future grain yield trends and providing a reference for research on grain yield prediction methods in Anhui Province.

2. Data sources and data preprocessing

2.1. Data sources

This study selected grain yield data of Anhui Province from 1989 to 2023, along with a set of relevant indicator data, which were compiled from the China Statistical Yearbook and the Anhui Statistical Yearbook. Meanwhile, using the Google Earth Engine (GEE) platform, we obtained the Normalized Difference Vegetation Index (NDVI) data from the MODIS Terra satellite and meteorological data within the boundaries of Anhui Province for the same time period.

2.2. Data preprocessing

First, the data were cleaned by removing invalid records and those with missing key fields, and the remote sensing data and yield data were temporally aligned to ensure consistency in their time ranges. Second, from the monthly remote sensing data, we extracted 16 statistical features for the growing season (April–September), including NDVI integral, NDVI peak, cumulative sunshine duration, etc., converting the monthly timeseries data into annuallevel features to align the temporal granularity with the annual data. Finally, since the features differ substantially in their scales (e.g., agricultural power generation is on the order of 10^5 , while agricultural electricity consumption is on the order of 10^2 , with a difference exceeding 1000fold), we applied Zscore standardization to all features. This ensured that each feature carried equal

weight in model training and prevented scale disparities from affecting model performance.

2.3. Feature selection

A twostep procedure was adopted to screen the seventeen annual indicators, aiming to eliminate multicollinearity and select the most predictive features.

2.3.1. Step 1: Collinearity removal

Pearson correlation coefficients were calculated among all features, with a correlation threshold set at 0.95. When the correlation between two features exceeded this threshold, the feature with the lower correlation was removed. A total of 21 highly correlated feature pairs were identified, leading to the removal of 8 redundant features and retaining 9 features.

2.3.2. Step 2: Random forest feature importance validation

A Random Forest model was used to evaluate the importance of each feature. The results showed that irrigated area and sown area were the two most important features, with importance values of 0.2810 and 0.2652, respectively. This is consistent with agricultural production knowledge—irrigation conditions and sown area are the most critical factors affecting grain yield^[16].

2.3.3. Final selection results

After the twostep screening, 8 annual key features were determined as model inputs, namely: irrigated area, sown area, agricultural electricity consumption, agricultural power generation, damaged area, pure nutrient equivalent of chemical fertilizer, affected area, and soil erosion area.

3. Model construction

3.1. XGBoost model

XGBoost (Extreme Gradient Boosting) is an ensemble learning algorithm based on Gradient Boosting Decision Tree (GBDT), proposed by Chen and Guestrin in 2016. It iteratively trains multiple decision trees to minimize the loss function. Compared with traditional GBDT, XGBoost incorporates a regularization term into the loss function to control tree complexity and prevent overfitting, while also introducing secondorder derivative information, enabling faster convergence and higher accuracy.

The core idea of XGBoost is to train multiple decision trees sequentially, where each tree fits the residuals of the previous model, and the final prediction is obtained by the weighted sum of the predictions from all trees. This “stepwise approximation” strategy allows the model to continuously correct previous prediction errors, thereby achieving high predictive accuracy^[17].

The objective function of XGBoost consists of a loss function and a regularization term:

$$L(\phi) = \sum_{i=1}^n l(y_i, \hat{y}_i^{(t)}) + \sum_{k=1}^t \Omega(f_k)$$

where $l(y_i, \hat{y}_i^{(t)})$ is the loss function of the t -th round, and $\Omega(f_k)$ is the regularization term of the k -th tree, which is used to control the complexity of the model:

$$\Omega(f_k) = \gamma T + \frac{1}{2} \lambda \sum_{j=1}^T w_j^2$$

where T is the number of leaf nodes in the tree, w_j^2 is the weight of the j -th leaf node, γ controls the complexity of leaf nodes, and λ controls the degree of regularization of the weights. By adjusting these parameters, a balance can be achieved between the model's fitting ability and its generalization capability.

3.2. MLP neural network model

MLP (MultiLayer Perceptron) is a feedforward neural network consisting of an input layer, one or more hidden layers, and an output layer. Each neuron in this network is connected to neurons in the preceding layer via weights, and the input signals, after being processed by a nonlinear activation function, are passed to the next layer. MLP is capable of learning complex nonlinear relationships among features and is also applicable to data with temporal characteristics.

The forward propagation process of MLP is as follows:

$$z_1 = W_1 x + b_1$$

$$a_1 = \sigma(z_1)$$

$$z_2 = W_2 a_1 + b_2$$

$$a_2 = \sigma(z_2)$$

$$\hat{y} = W_3 a_2 + b_3$$

Among them, W_i and b_i are the weight matrix and bias vector of the i -th layer, respectively, and $\sigma(\cdot)$ denotes the activation function (ReLU is used in this paper). By stacking multiple layers of nonlinear transformations, the MLP is able to learn deep representations of the input features, thereby capturing complex nonlinear patterns.

The MLP is trained using the backpropagation algorithm, which employs gradient descent to minimize the mean squared error between the predicted values and the ground truth. During the training process, the error signal is propagated backward layer by layer from the output layer to the input layer, and the weight and bias parameters of each layer are updated accordingly, driving the model to approach the optimal solution gradually.

3.3. Random forest model

Random Forest, proposed by Breiman in 2001, is an ensemble learning algorithm based on decision trees. It trains a large number of decision trees in a parallel manner, where each tree is built using random sampling of both samples and features. The final prediction is produced by voting (for classification) or averaging (for regression). This "collective decision" mechanism reduces the variance of the model and thus improves its generalization ability.

The prediction process of Random Forest is as follows:

$$\hat{y} = \frac{1}{K} \sum_{k=1}^K f_k(x)$$

Among them, K represents the number of decision trees, and $f_k(x)$ denotes the predicted value of

the k -th tree. By averaging the prediction results from multiple trees, the random forest can counteract the fluctuations caused by the randomness of a single tree, thereby yielding more stable predictions.

3.4. Stacking ensemble framework

Stacking (stacked generalization) is a hierarchical ensemble learning method proposed by Wolpert in 1992, which employs a metalearner to learn from the prediction results of base models, thereby achieving adaptive fusion of different models. Unlike simple weighted averaging or voting strategies, Stacking automatically learns the optimal weights of each base model by training a metalearner, making better use of the complementary information across different models.

The fundamental idea of Stacking is to transform the model fusion problem into a new learning problem: using the prediction outputs of base models as input features and the original labels as the target, the metalearner is trained to learn how to combine these predictions optimally. This strategy adaptively adjusts the weights of each base model according to the data characteristics, without requiring manual specification.

3.5. Metalearner and ensemble strategy

In this study, Ridge regression is adopted as the metalearner, which learns the weights of the base models by minimizing the regularized squared error:

$$\min_w \sum_{i=1}^n \left(y_i - \sum_{m=1}^M w_m \hat{y}_{i,m} \right)^2 + \alpha \sum_{m=1}^M w_m^2$$

where w_m denotes the weight of the m base model, α is the regularization coefficient (set to 1.0 in this study), and M is the number of base models ($M=3$ in this study).

L2 regularization is introduced in Ridge regression to constrain the magnitude of the weights, preventing any single base model from being assigned an excessively large weight, thereby ensuring the stability of the ensemble model. Moreover, Ridge regression has a closedform solution, which allows for relatively fast training. See **Figure 1**.

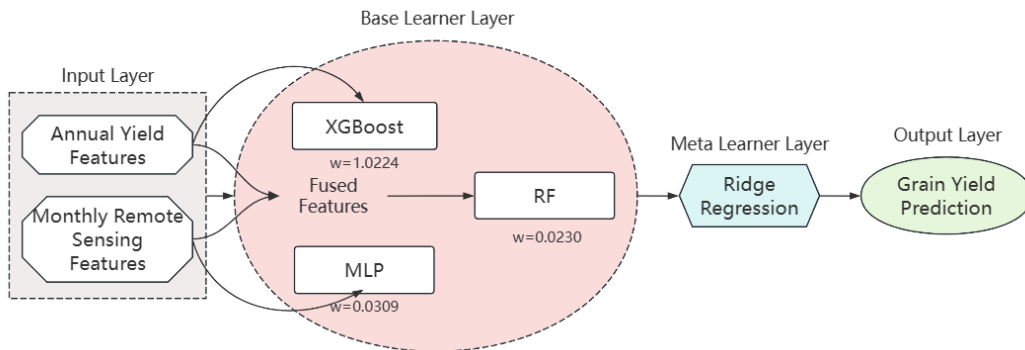


Figure 1. Model structure.

4. Model prediction

4.1. Experimental design

4.1.1. Data partitioning and validation strategy

In this study, leaveoneout crossvalidation (LOOCV) was adopted to evaluate model performance. Specifically, data from one year were selected as the test set, while the data from the remaining 34 years

were used as the training set; this procedure was repeated 35 times, covering all years from 1989 to 2023. This approach makes full use of the limited sample and avoids the fluctuations in results that may arise from random partitioning, thereby ensuring more reliable evaluation outcomes.

4.1.2. Evaluation metrics

Three metrics were selected to assess model performance:

- (1) Coefficient of determination (R^2)
Reflects the proportion of variance explained by the model; the closer the value is to 1, the better the goodness of fit.
- (2) Root mean square error (RMSE)
Represents the standard deviation of prediction errors, expressed in kg/ha.
- (3) Mean absolute percentage error (MAPE)
Reflects the average level of relative error, presented as a percentage.

4.1.3. Baseline model settings

To examine the effectiveness of the Stacking ensemble model, three groups of baseline models were established:

- (1) XGBoost baseline model: trained using only annual structured features.
- (2) Direct feature fusion model: annual and monthly features were directly concatenated and then fed into XGBoost for training.
- (3) Stacking ensemble model: integrates the prediction results from three base models, namely XGBoost, MLP, and RandomForest.

4.2. Model performance comparison

4.2.1. Overall performance evaluation

It can be seen from Table 1 that the Stacking ensemble model outperforms the other models in terms of both R^2 and RMSE:

- (1) The R^2 reaches 0.9575, which is an improvement of 0.0039 over the XGBoost baseline model;
- (2) The RMSE is 140.90 kg/ha, which is 6.21 lower than that of the XGBoost baseline model;
- (3) The MAPE is 4.08%, which is comparable to that of the XGBoost baseline model. This indicates that the Stacking ensemble strategy can effectively integrate multi-source information and improve prediction accuracy.

Table 1. Performance comparison of different models

Model	R^2	RMSE	MAPE
XGBoost baseline	0.9536	147.11	4.07%
Direct feature fusion model	0.9296	181.23	4.51%
Stacking ensemble model	0.9575	140.9	4.08%

4.2.2. Comparison between actual and predicted values

Figure 2 presents scatter plots of actual vs. predicted values for the Stacking ensemble model and the XGBoost baseline model. In the figure, the diagonal line represents perfect prediction, and the closer the

scatter points are to the diagonal, the higher the prediction accuracy.

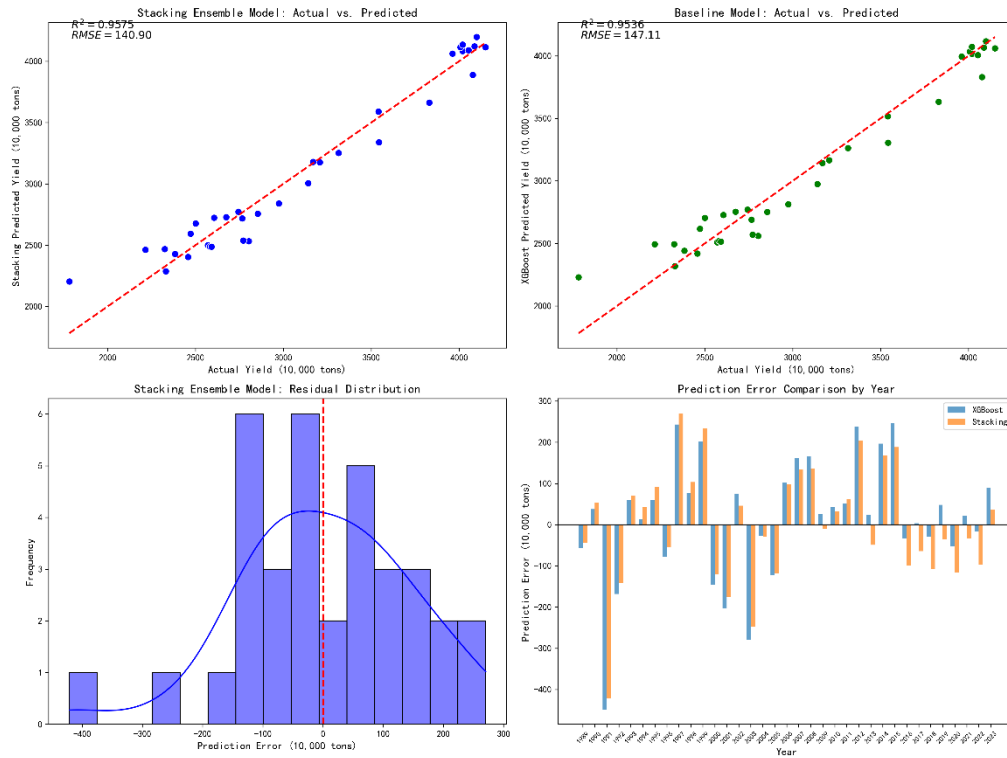


Figure 2. Comparison of model prediction results.

As can be seen from **Figure 2**:

- (1) The upper left subplot shows the prediction results of the Stacking ensemble model, with scatter points concentrated near the diagonal line, $R^2 = 0.9575$, and $RMSE = 1.4090$ million tons, indicating relatively high prediction accuracy;
- (2) The upper right subplot shows the XGBoost baseline model, with scatter points relatively more dispersed, $R^2 = 0.9536$, and $RMSE = 1.4711$ million tons, indicating slightly lower accuracy than the Stacking model;
- (3) The lower left subplot shows the residual distribution; the residuals of the Stacking model approximately follow a normal distribution with a mean close to zero, indicating little model bias;
- (4) The lower right subplot shows the error comparison; in most years, the absolute prediction error of the Stacking model is smaller than that of the XGBoost baseline model, especially in years with larger errors such as 1991 and 2003, where the Stacking model demonstrates a more pronounced effect in reducing errors.

4.2.3. Comparison of time series prediction

Figure 3 presents a comparison of the prediction curves generated by the three models:

- (1) In terms of overall trend, all three curves are able to track the variations in actual yield, particularly during the period of steady growth after 2012;
- (2) The Stacking model demonstrates a clear advantage, with its prediction curve more closely matching the

- actual yield curve, and it also captures yield fluctuations more accurately in years with large variations, such as 1991 and 2003;
- (3) The direct feature fusion model performs relatively poorly, with larger prediction errors and notably overestimated values in lowyield years, indicating that simply combining features does not improve model performance.

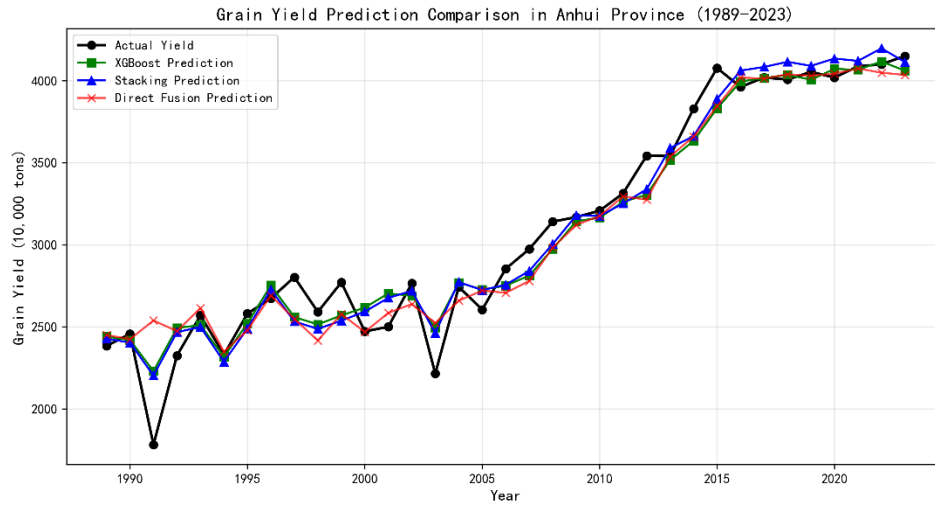


Figure 3. Comparison of timeseries forecasting of grain yield in Anhui Province (1989–2023).

4.2.4. P erformance improvement analysis

Table 2. Performance improvement of the Stacking ensemble model over the baseline models

Indicator	Baseline model	Stacking model	Improvement
R^2	0.9536	0.9575	0.41%
RMSE	147.11	140.90	−4.22%
MAPE	4.07%	4.08%	0.25%

Although the improvements in R^2 and RMSE are not substantial, in the context of grain yield prediction with a limited sample size, even marginal performance gains still hold practical value. The RMSE decreased by 6.21, corresponding to an improvement in prediction accuracy of 4.22%. As Anhui Province is a major grainproducing region in China, the accuracy requirements for yield forecasting are inherently high, and these gains are therefore valuable in realworld applications.

4.2.5. Comparison with direct feature fusion

The direct feature fusion model performed even worse than the baseline model, with R^2 decreasing by 0.0240 and RMSE increasing by 34.12. This indicates that simply concatenating annual features and monthly features does not improve predictive performance; rather, it may lead to overfitting due to increased feature dimensionality and mutual interference among features. In contrast, the Stacking ensemble strategy, which automatically learns the weights of each base model through the metalearner, is better able to exploit the complementary information across different feature types.

5. Results and discussion

5.1. Value and significance of multisource data fusion

Grain yield prediction is a highly complex task involving multiple interacting factors from the natural, social, and economic spheres. Relying on a single data source is insufficient to comprehensively cover these factors, making multisource data fusion a key approach to improving prediction accuracy. In this study, by employing the Stacking ensemble learning framework to integrate annual structural features with monthly timeseries features, we validated the effectiveness of multisource data fusion in grain yield prediction.

In terms of model performance, the Stacking ensemble model that fused annual and monthly data achieved an R^2 of 0.9575, which is an improvement of 0.0039 over the XGBoost baseline model that used only annual features, and the RMSE decreased by 6.21. Although the improvement is modest, it remains practically meaningful for smallsample regression problems such as grain yield prediction. An RMSE reduction of 6.21 corresponds to a 4.22% improvement in prediction accuracy. For Anhui Province, as a major grainproducing region, this level of accuracy enhancement can provide more reliable decisionmaking support for government authorities.

With regard to feature contributions, annual features and monthly features play distinct roles in yield prediction. Among the annual features, sown area and cultivated area are the most critical predictors, together accounting for 60.4% of feature importance, indicating that the fundamental input factors of agricultural production exert a decisive influence on yield. Although the weights of monthly features such as NDVI and sunshine duration are relatively small (approximately 5%), they provide growthseason dynamic information that annual features cannot capture, and in certain years they notably improve prediction accuracy. This demonstrates that multisource data fusion not only enhances overall prediction accuracy but also captures subtle variations that are difficult to reflect using a single data source.

Regarding the ensemble strategy, the Stacking ensemble model effectively fuses multisource information by adaptively learning the weights of each base model through the metalearner. XGBoost, as the primary contributor (weight of 1.0224), captures annualscale structural information; MLP and RandomForest, as secondary contributors (weights of 0.0309 and 0.0230, respectively), capture timeseries variation patterns within the growing season and interactions among features. The weights of all three base models are positive, indicating that their predictions are complementary, thereby verifying the effectiveness of the Stacking ensemble strategy.

The significance of multisource data fusion in grain yield prediction lies not only in improving prediction accuracy but also in providing technical support for the development of intelligent agriculture. By integrating multisource data—including remote sensing, meteorology, soil, and crop information—and building intelligent agricultural decisionmaking systems, it is possible to achieve precise, automated, and intelligent agricultural production. This holds strategic importance for addressing climate change challenges, safeguarding national food security, and promoting sustainable agricultural development.

5.2. Practical applications and policy implications of the model

The grain yield prediction model constructed in this study has practical application value and can provide reference for government decisionmaking and sustainable agricultural development.

At the level of government decisionmaking, accurate yield predictions enable authorities to grasp the grain production situation in a timely manner and formulate appropriate grain policies. If the predictions indicate a potential decline in yield, the government can intervene proactively—for example, by increasing

agricultural subsidies, promoting highyield technologies, or strengthening disaster prevention and control—to stabilize grain production. The model itself can also identify key factors affecting yield, providing a reference for policy formulation. For instance, since sown area and cultivated area are core influencing factors, this suggests that governments should prioritize infrastructure construction such as farmland water conservancy and land consolidation, while ensuring that the quantity of cultivated land is maintained and its quality is improved.

For farmers, the feature importance results output by the model can help them understand which factors have the greatest impact on yield, allowing them to adjust their agricultural input structure accordingly and improve resource use efficiency. Fertilizer is one of the important factors affecting yield; based on the analytical results, farmers can apply fertilizer more rationally, promote precision fertilization techniques, and reduce waste and production costs.

In the field of precision agriculture, the application of remote sensing data plays a critical supporting role. By monitoring crop growth dynamics in real time, problems can be detected and addressed promptly, enhancing the precision of agricultural management. The NDVI feature reflects realtime crop growth status; when its value falls below the normal threshold for the corresponding period, farmers can implement timely field management measures such as irrigation and fertilization to help restore crop growth.

5.3. Research limitations and future directions

This study has certain limitations that need to be addressed in future work.

In terms of data, the limited sample size is a major shortcoming. This study used only 35 years of data from 1989 to 2023, and the relatively small sample size constrains the model's generalization ability to some extent. NDVI data for certain years also suffered from missing values or quality issues, which affected the accuracy of feature extraction. Future work could collect longer timeseries data to expand the sample and strengthen data preprocessing procedures to improve overall data quality.

With regard to the model, the relatively high complexity is another issue. The Stacking ensemble model, comprising three base models and one metalearner, has a deep architecture that entails increasing computational costs. The interpretability of the model is also somewhat limited, as its internal decisionmaking process remains difficult to understand intuitively. Future research could explore more concise ensemble strategies to reduce complexity and introduce interpretability tools such as SHAP values and LIME to deconstruct the decision logic and enhance model transparency quantitatively.

In terms of application, prediction errors of the model are notably larger in years with natural disasters, primarily because such extreme events account for a very small proportion of the training samples. Future work could supplement the dataset with historical disasteryear data to construct prediction modules tailored to extreme scenarios, thereby strengthening the model's ability to handle anomalous situations. In addition, this model was trained and validated exclusively on data from Anhui Province; direct transfer to other provinces or regions may not be directly applicable, and parameter recalibration and regional validation would be required.

In summary, grain yield prediction is a continuously evolving research field. With increasing data volumes, advances in modeling techniques, and changing application requirements, future research will move toward higher accuracy, stronger interpretability, and broader applicability, thereby making greater contributions to ensuring national food security.

Disclosure statement

The authors declare no conflict of interest.

References

- [1] Gao X, Dong Y, Xu W, et al., 2025, Analysis and Prediction of Spatiotemporal Changes in Grain Production in Central Asia Based on ARIMA Model. *Journal of University of Chinese Academy of Sciences*, 42(4): 472–486. (in Chinese)
- [2] Zhang W, Sun D, Wang Y, et al., 2019, Prediction of Grain Yield in Liaoning Province Based on Support Vector Machine. *Economic Mathematics*, 36(1): 96–99. (in Chinese)
- [3] Dong M, Miao L, 2025, Research on Grain Yield Prediction Based on ARIMA-X-RFR Model. *Journal of Jiamusi University (Natural Science Edition)*, 43(11): 21–24. (in Chinese)
- [4] Jia M, Cai Z, Hu J, et al., 2021, Research on Grain Yield Prediction Model Based on Machine Learning. *Journal of Hebei Agricultural University*, 44(3): 103–108. (in Chinese)
- [5] Shen Y, Mercatoris B, Cao Z, et al., 2022, Improving Wheat Yield Prediction Accuracy Using LSTM-RF Framework Based on UAV Thermal Infrared and Multispectral Imagery. *Agriculture*, 12(6): 892.
- [6] Wang X, Huang J, Feng Q, et al., 2020, Winter Wheat Yield Prediction at County Level and Uncertainty Analysis in Main Wheat-Producing Regions of China with Deep Learning Approaches. *Remote Sensing*, 12(11): 1744.
- [7] Zhao Y, Fan L, Qin Z, et al., 2025, Analysis of Influencing Factors and Combined Prediction of Grain Yield Based on SVR and CNN-LSTM-ATTENTION Model. *Journal of the Chinese Cereals and Oils Association*: 1–12. (in Chinese)
- [8] Ma D, Zhao J, Yan W, et al., 2025, Grain Yield Prediction in Southern China Based on Stacking Ensemble Algorithm. *Hubei Agricultural Sciences*, 64(5): 155–159 + 184. (in Chinese)
- [9] Zou Y, 2024, Research on Maize Yield Prediction in China Based on Machine Learning, thesis, Nanjing University of Information Science and Technology (in Chinese)
- [10] Basso B, Ritchie J, Pierce F, et al., 2011, Spatial Validation of Crop Models for Precision Agriculture. *Agricultural Systems*, 68: 97–112.
- [11] Liu J, He X, Wang P, et al., 2019, Early Prediction of Winter Wheat Yield Using Long Time Series Meteorological Data Combined with Random Forest Method. *Transactions of the Chinese Society of Agricultural Engineering*, 35(6): 166–174. (in Chinese)
- [12] Panov D, Sakharova E, 2022, Using Radar Data for Grain Crops Yield Forecasting in the Novosibirsk Region. *Russian Meteorology and Hydrology*, 47(6): 473–478.
- [13] Zhang L, 2025, Research on Grain Yield Prediction Model Based on Deep Learning, thesis, Changchun University of Technology (in Chinese)
- [14] Jeong J, Resop J, Mueller N, et al., 2016, Random Forests for Global and Regional Crop Yield Predictions. *PLOS ONE*, 11(6): 56–57.
- [15] Brown J, Hochman Z, Holzworth D, et al., 2018, Seasonal Climate Forecasts Provide More Definitive and Accurate Crop Yield Predictions. *Agricultural and Forest Meteorology*, 260: 247–254.
- [16] Chen C, Lin Z, Sun X, 2007, Dynamic Prediction and Suggestions for Grain Yield in Shandong Province. *Journal of Anhui Agricultural Sciences*, 35(28): 8773–8775. (in Chinese)
- [17] Tian Q, Li J, Ding Y, et al., 2025, Machine Learning-Driven Grain Yield Prediction and Influencing Factor Analysis in Heilongjiang Province. *Grain and Feed Science and Technology*, 2025(2): 13–15. (in Chinese)

Publisher's note

Bio-Byword Scientific Publishing remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.