

Forecasting cropping patterns to increase crop yields food and horticulture using a machine approach learning

Imam Sanjaya^{1*}, *Teddy Mantoro*¹, *Jelita Asian*¹, *Ivana Lucia Kharisma*¹, and *Muhammad Ikhsan Thohir*¹

¹Departement of Informatic Engineering, Nusa Putra University, Sukabumi, West Java, Indonesia

Abstract. Nearly 82% of rural areas still rely on the agricultural sector, one of which is the cultivation of food crops and horticulture. Changing climatic conditions are one of the causes of farmers' failure to predict the selection of the right time for the cultivation process. This work predicts the selection of the right time for the cultivation of certain crops in order to optimize crop yields. Forecasting uses several machine learning methods by comparing the best results. The results showed that machine learning could produce good information at the right time and on certain types of plants to be cultivated in an area. Predictive recommendations for planting corn in 2022 are optimum planting in January to March 2022 and in August to December 2022. The optimum fertilizer application is N fertilizer dose = 100-270 kg/ha, P fertilizer dose = 63-100 kg/ha, and K fertilizer dose = 156-200 kg/ha.

1 Introduction

Food availability is a basic human need for survival [1]. Based on physiological and anatomical traits, each group of plants has its unique set of characteristics, resulting in variations in harvest results. Food shortages are common in a number of nations, including Indonesia. One of the causes is crop failure. Most of these crop failures are caused by environmental causes, specifically climate change. Rainfall, temperature, and air humidity all have an impact on the climate [2]. Most agricultural activities are dictated by the balance between the amount of water available on the land and the amount of water required for its growth [3].

Water availability influences planting patterns in a given area [4]. Planting patterns are developed to discover when plants can most efficiently evaluate water availability [5]. One of the planting pattern activities is picking plants whose water needs agree with the distribution of rainfall, as well as crop rotation that can efficiently utilize rainfall.

One of the plants that must be correct in defining the planting pattern is the food crop maize because this corn has a distinctive water demand that must be adjusted to both the vegetative and generative stages [6]. Besides being used by humans as a substitute for rice

* Corresponding author: ivana.lucia@nusaputra.ac.id

[7], corn is also used as poultry feed [8]. Corn for chicken feed must still be imported since the domestic corn supply is insufficient to meet demand.

One of the reasons for simultaneous corn planting by Indonesian farmers is a lack of knowledge and use of computer information technology, particularly Machine Learning, which can forecast or predict planting patterns based on rainwater availability, which can influence fertilizer and nutrient absorption by plants.

The national productivity of corn commodities in Indonesia continues to increase every year. In 2012, corn productivity reached 4.5 tons/ha-1 and then experienced successive increases in 2013-2016, namely 4.84; 4.95; 5.18; and 5.31 tons/ha-1. However, the production and productivity of corn plants in Sukabumi Regency fluctuates every year. Corn productivity in Sukabumi Regency in 2012 was 5.5 tons/ha-1, then in 2013 it decreased to 5.4 tons/ha-1 and in 2014 it increased to 5.8 tons/ha-1 [9]. Climate change from global warming is regarded to be one of the causes of instability in corn productivity in Indonesia. Global warming affects the availability of water in rainfed areas. Corn plants require moderate water availability, but water availability that is too little or too much will result in reduced corn output [9].

In this research, efforts were made to process rainfall data to determine the best planting pattern for corn plants using several machine learning methods, namely the k-NN, Naïve Bayes, and Decision Tree methods [10] The KNN method was chosen because of its capacity to categorize the unknown effects of fertilizer use and water availability on maize plants using training and test data. KNN can generate mathematical algorithms to convert the values of these criteria into classification statements. This approach accurately classifies data by selecting the nearest neighbor K value correctly [11].

Crop forecasting and monitoring rely primarily on meteorological data, which are not always available and often have low regional representativeness. Using the system created and maintained by the European Commission's Joint Research Centre, along with its gridded daily meteorological observations, biased-corrected reanalysis AgMERRA, and the ERA-Interim reanalysis, it evaluates the viability of reanalysis-based crop monitoring and forecasting in this study [12]. Studies on single cropping reveal significant variance in sensor utilization. Historically, the most widely used sensors for mapping crop varieties within fields were optical sensors. High and extremely high spatial resolution optical sensors (like UAVs) have made it possible to precise identification of a field's single cropping [12] ERS, SkyMed, RADARSAT-2, Sentinel-1, and other high temporal resolution microwave sensors have all been used to map single crops both separately and in conjunction with optical sensors [13–15].

The k-Nearest Neighbor algorithm is a supervised learning technique that uses the majority of k-nearest neighbor categories to categorize new sample results. The purpose of this method is to categorize new objects based on their properties and training data samples. The learning data closest to the object is utilized to classify neighborhoods using the k-Nearest Neighbor method. The learning data is projected into a multidimensional space, with each dimension representing a feature of the data [16]. A straightforward probabilistic classifier, Naïve Bayes adds up the frequencies and value combinations from a dataset to determine a set of probabilities [17]. One popular and highly effective technique for categorization and prediction is the use of decision trees. Using the decision tree method, a decision tree representing rules can be created from very big data [18].

2 Material and Methods

At the first stage of research, its task is to review papers relevant to the topic [19]. Information is gathered using terms like "the influence of water on nutrient availability," "maize crop planting patterns," "the use of machine learning for planting patterns," and "the

prediction of planting patterns using machine learning." All of these search terms are used to search a variety of databases, including papers and journals. After gathering papers, the literature analysis focuses on predicting corn planting patterns using machine learning, which investigates water availability in soil. At this stage, data is also collected from the Department of Agriculture in the form of data on soil fertility conditions in the form of pH and main nutrients, namely Nitrogen (N), Phosphorus (P) and Potassium (K). These nutrients will be influenced by the availability of water so that they can be absorbed by corn plants. From BMKG or BPS in the form of climate condition data in the form of rainfall (wet months – humid months – dry months). This climate data will provide an overview of rainfall every month.

At second stage, testing the classification model using the k-NN, naïve Bayes, decision tree, LSTM approach from a collection of climate condition data (will provide information on the availability of water in the soil which helps nutrients become available as plant nutrition) [20]. The results of this classification are expected to provide information regarding the appropriate planting pattern for corn plants to grow optimally.

At third stage, it is carried out using two proposed methods. The first process is a qualitative study, using existing data from other studies that are significant to this research [21]. The second process is a quantitative study. At this stage we will design predictions of planting patterns from the classification process using the k-NN, naïve Bayes and decision tree approaches for corn plants. This study process will be sourced from extracting information in the form of climate data in Sukabumi Regency. The data taken is in the form of rainfall data and rainfall categories. This climate data will be tested using machine learning to determine the optimum rainwater requirements for dissolving N, P and K fertilizers which affect the growth of corn plants. In a quantitative approach to classification design, experiments will be carried out to measure the availability [22] of N, P, and K fertilizer in wet and dry soil conditions. These wet and dry soil characteristics reflect the climate conditions. Plants can absorb nutrients in proportion to the amount of water they absorb and the concentration of nutrients in the water [23]. So, if the soil is dry, the roots are unable to absorb the nutrients found in the water. Rainfall conditions are classified as wet, humid, or dry based on the volume of rainfall that reaches the ground each year.

3 Results and Discussion

3.1 Research Implementation

This research was carried out in one of the agricultural centers in Sukabumi Regency, namely in Nagrak District. There are 3 different fields and corn planting times. Condition 1 was observed from 04 October to 23 December 2020. Condition 2 was observed from 10 December 2020 to 06 March 2021. Condition 3 was observed from 20 March to 07 June 2021. Observations of the three conditions were carried out on different land and different times, this was done to see the response of rainfall, temperature and fertilizer to corn yields. Data collection was carried out at several points on the same land and the yield values were averaged. Temperature measurements were taken in the morning and evening. Fertilizer doses of Nitrogen (N), Phosphate (P), and Potassium (K) are adjusted to the land conditions and needs of each corn plant. Rainfall measurements were taken from the Mathematics and Geophysics Agency (BMKG) in West Java. Rain categories are divided into low rainfall 0 – 100mm (1), medium rainfall 101 – 250mm (2), and high rainfall 251 – 500mm (3). The data is divided into 80 percent (80%) training data and 20 percent (20%) test data [24]. The test was carried out by comparing 3 algorithms, namely k-NN, Naïve Bayes and decision tree to

determine the best prediction value for planting time on the land. The following Table 1 shows samples of data collected.

Table 1. Data on Rainfall Samples, N, P, K Fertilizer Doses, and Harvest Yields

Month	Average Rain	Rain Category	Average Temperature	N (Kg/Ha)	P (Kg/Ha)	K (Kg/Ha)	Plant Result (kg/1000m2)
100	191.8	2	23.8	154	63	184	600
100	191.8	2	22.5	154	63	184	600
100	191.8	2	23.4	154	63	184	600
100	191.8	2	24	154	63	184	600
100	191.8	2	23.6	154	63	184	600

Table 1 depicts a total of 248 data, consisting of 8 columns, 4 columns as attributes (Month, Average Rain, Rain Category, Average Temperature) and 3 columns as labels (N, P, and K). The parameter variables for this experiment are carried out in Table 2.

Table 2. Plant Yield Measurement Parameters

Description	Variables
Month	Corn planting time components (as attributes)
AVERAGE rain	Daily rainfall conditions taken from BMKG (as attribute)
Rain category	Differences in rainfall from low, medium, and high (as attributes)
Average temperature	Horizontal coordinate point access point distance (in meters) from object/user (as attribute)
N	Condition of Nitrogen fertilizer applied to corn fields (as label)
P	State of Phosphate fertilizer applied to corn fields (as label)
K	Condition of Potassium fertilizer applied to corn fields (as label)

An independent variable is a variable that exists independently or that influences or causes changes or the appearance of the dependent variable. In this study, the independent variables include month, average rainfall, rain category, and average temperature. Meanwhile, the dependent variable is a variable that is influenced or results from the presence of an independent variable [25]. In this study, the dependent variables are N, P, and K. The independent variable in this study will explain and influence the dependent variable through water availability, which will affect the absorption of N, P, and K fertilizer by corn plant roots in the soil. Systematic explanations for phenomena.

3.2 Data Splitting

Python's Sckit-Learn (sklearn) can be used to easily implement train/test splits. To be able to use it, we need to import Sklearn first. After that we can use the train_test_split() function. Source code given in Figure 1.

```
from sklearn.model_selection import train_test_split
```

Fig. 1. Source Code Data Split

After that, it continues with the definition of the source data (x) and the target data (y). The source data is the month column, Average rainfall, rain category, average temperature, N, P, and K. While the target data is the crop yield column. Then we can implement train/test. Source code given in Figure 2.

```
[15] x_train, x_test, y_train, y_test = train_test_split(x, y, test_size = 0.2, random_state = 123)
```

Fig. 2. Source Code Train Test Split

- x_train: To accommodate the source data to be trained
- x_test: to hold the target data to be trained
 - y_train: to hold source data that will be used for testing
 - y_test: to hold target data that will be used for testing

x and y are the variable names that will be used when defining the source data and target data. The test_size parameter is used to define the size of the testing data. In Figure 2 test_size=0.2 means the data used for testing is 20% of the entire dataset.

Next, the variable x will be scaled. There are three scalers in the scikit-learn library that are often used, namely: StandardScaler, MinMaxScaler, and RobustScaler. Used in the preprocessing process. In Figure 3, StandardScaler is used, this scaler removes the mean (centered at 0) and scales to the variance (standard deviation = 1), assuming the data is normally distributed (Gauss) for all features.

3.3 Model Evaluation

Confusion matrix (error matrix) is a predictive analytical tool for displaying and comparing actual values with model predicted values which can be used to produce evaluation matrices such as accuracy, precision, recall and F1-Score (F-Measure) [26]. In other words, the confusion matrix is described as a table that is often used to measure the performance of machine learning classification models. Predicted model using machine learning algorithm will be evaluated its performance , showed in Confusion Matrix and Classification report. Confusion matrix and classification report using KNN-3 models decribed in Figure 3.

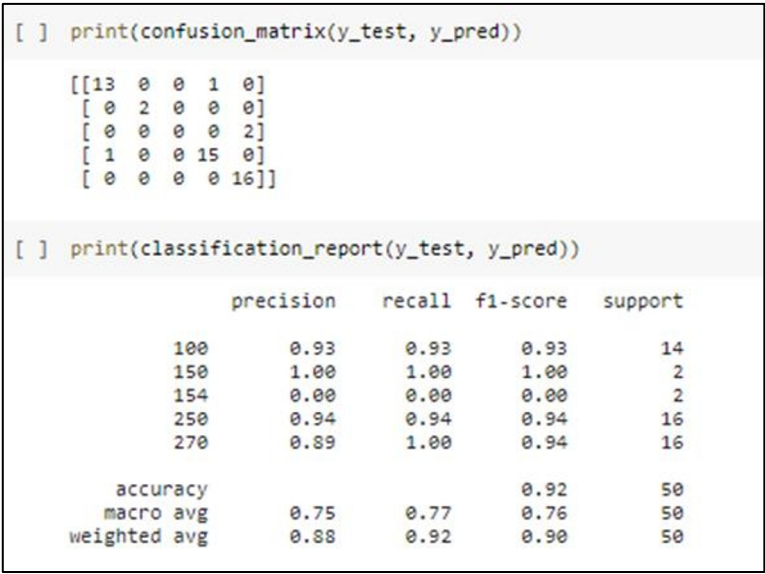


Fig. 3. Model evaluation using KNN-3

From the results of the confusion matrix above, the ratio of TP (true positive), FP (false positive), FN (false negative), and TN (true negative) can be calculated as follows:

- $TP = (13+2+0+15+16) = 46$ (Corn yield data is good and the model also provides good yield predictions)
- $FP = (0+0+1+0) = 1$ (Maize yield data is not good but the model provides good yield predictions)
- $FN = (0+3+1+0) = 4$ (The corn yield data is good but the model predicts the results are not good)
- $TN = (2+2+15+16) = 35$ (Maize yield data is not good and the model predicts the results are not good)
- Accuracy Value: 0.92 (92%) (Is the ratio of correct predictions (positive and negative) to the entire data, thus accuracy is the level of closeness of the predicted value to the actual value).
- Precision : $(0.93+1+0+0.94+0.89)/5 = 0.75$ (75%)
- Recall: $(0.93+1+0+0.94+1)/5 = 0.77$ (77%)
- F1-score : $(0.93+1+0+0.94+0.94)/5 = 0.76$ (76%).

Confusion matrix and classification report using Naïve Bayes models decribed in Figure

4.

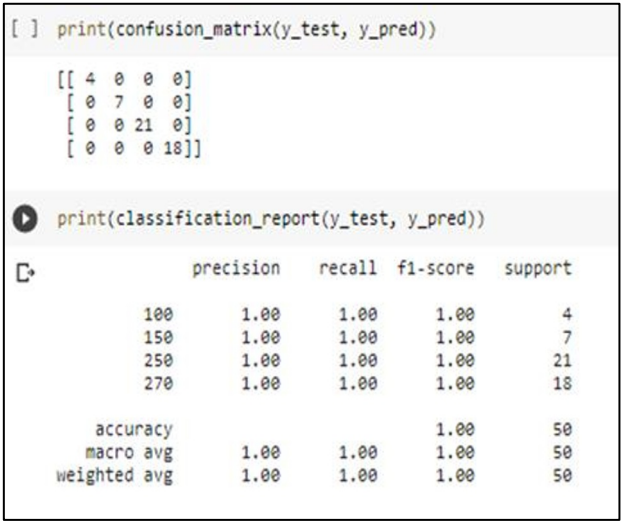


Fig. 4. Model evaluation using Naïve Bayes

- $TP = (4+7+21+18) = 50$ (Maize yield data is good and the model also provides good yield predictions)
- $FP = (0+0+0+0) = 0$ (Maize yield data is not good but the model provides good yield predictions)
- $FN = (0+0+0+0) = 0$ (The corn yield data is good but the model predicts the results are not good)
- $TN = (7+21++18) = 46$ (Maize yield data is not good and the model predicts the results are not good)
- Accuracy Value: 1.00 (100%) (Is the ratio of correct predictions (positive and negative) to the entire data, thus accuracy is the level of closeness of the predicted value to the actual value).
- Precision: $(1+1+1+1)/4 = 1$ (100%)
- Recall: $(1+1+1+1)/4 = 1$ (100%)
- F1-score : $(1+1+1+1)/4 = 1$ (100%)

Confusion matrix and classification report using Decision Tree models are described in Figure 5.

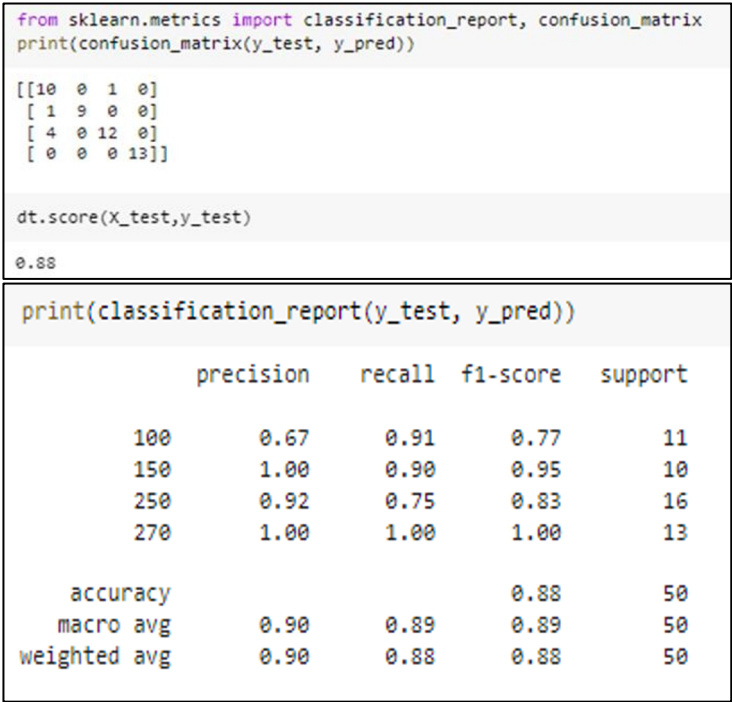


Fig. 5. Model evaluation using Decision Tree

From the results of the confusion matrix above, the ratio of TP (true positive), FP (false positive), FN (false negative), and TN (true negative) can be calculated as follows:

- TP = (10+9+12+13) = 44 (Maize yield data is good and the model also provides good yield predictions)
- FP = (1+1+4+0) = 6 (Maize yield data is not good but the model provides good yield predictions)
- FN = (5 +0+1+0) = 6 (The corn yield data is good but the model predicts the results are not good)
- TN = (9+12+13) = 34 (Maize yield data is not good and the model predicts the results are not good)
- Accuracy Value: 0.88 (88%) (Is the ratio of correct predictions (positive and negative) to the entire data, thus accuracy is the level of closeness of the predicted value to the actual value).
- Precision : (0.67+1+0.92+1)/4 = 0.90 (90%)
- Recall : (0.91+0.90+0.75+1)/4 = 0.89 (89%)
- F1-score : (0.77+0.95+0.83+1)/4 = 0.89 (89%)

The calculation results of FP and FN are close, meaning the accuracy value of this prediction is good.

Classification report using LSTM models decribed in Figure 6.

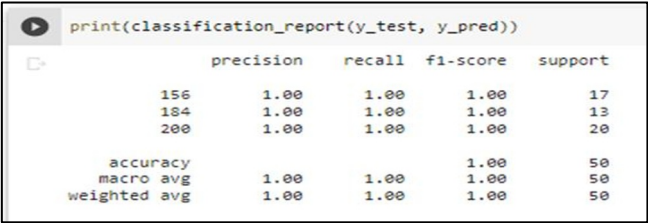


Fig. 6. Classification Report using LSTM

Based on Figure 6, the planting time classification results are displayed which show the magnitude of precision, recall, f1 score and support so that an accuracy result of 1 is obtained for the LSTM model. This provides a calculation that the model's accuracy value in classifying correctly is 100%. The overall value of precision is 1, this gives an idea of the accuracy of the requested data with the prediction results provided by the model being 100%. The overall value of recall is 1, this describes the success of the model in retrieving information by 100%.

3.4 Comparison of Classification Report Results for 4 Models

In this research, we can compare the classification reports of all the proposed models to obtain optimal results. From the confusion matrix calculation, optimum N fertilizer accuracy was obtained.

A comparison of classification reports for the 4 models is carried out in Table 3.

Table 3. Comparison of Classification Results of 4 Models

model	precision	recall	fi-Score	support	accuracy
k-NN 3	0.75	0.77	0.76	50	0.92
Naive bayes	1.00	1.00	1.00	50	1.00
Decision Tree	0.91	0.92	0.91	50	0.88
LSTM	1.00	1.00	1.00	50	1.00

Table 3 shows that Naïve Bayes and LSTM are the best models among the other models because they provide optimal recall, error rate and accuracy values. These results provide an illustration that the observed data is close to the predicted results. A comparison of fertilizer dosage reports can be seen in Table 4.

Table 4. Fertilier Dosage Results of 4 Models

Model	Fertilizer Dosage	precision	recall	fi-Score
k-NN3	100	0.93	0.93	0.93
	150	1.00	1.00	1.00
	250	0.94	0.94	0.94
	270	0.89	1.00	0.94
Naïve Bayes	100	1.00	1.00	1.00
	150	1.00	1.00	1.00
	250	1.00	1.00	1.00
	270	1.00	1.00	1.00
Decision Tree	100	0.67	0.91	0.77
	150	1.00	0.90	0.95
	250	0.92	0.75	0.83

	270	1.00	1.00	1.00
LSTM	100	1.00	1.00	1.00
	150	1.00	1.00	1.00
	250	1.00	1.00	1.00
	270	1.00	1.00	1.00

From Table 4, applying 100-270 kg/ha of fertilizer still produces good predictive values, this is because in this range of fertilizer application, corn plants still produce good harvests from field observations.

A comparison of research data and overall prediction results is shown in Figure 7 and 8.

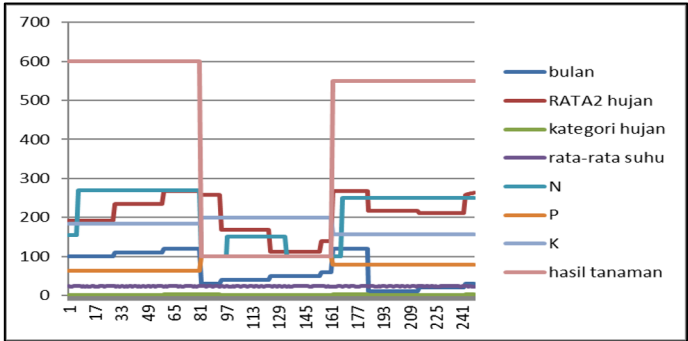


Fig. 7. Data from Corn Planting Research Results

The research results in Figure 7 show that the x line is the amount of data, and the y line is the result of the measurement value. The months of March – May provide less good harvest results even with reasonable fertilizer application, this is because in that month the research results show little rainfall.

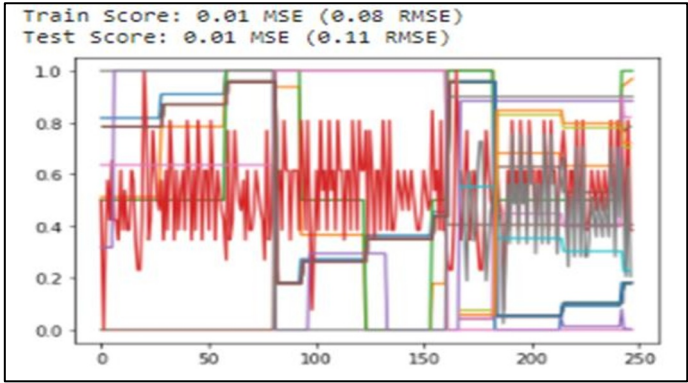


Fig. 8. Image of Prediction Results with LSMT

April – May, including rain category 1 (little rain) fertilizer dose = 100-150 kg /ha N, 100 kg/ha P, and 200 kg/ha K. The optimum crop yield condition is in the 25-60th data condition, namely rain category 2 (moderate rain), fertilizer dose = 270 kg/ha N, 63 kg/ha P, and 184 kg/ha K and in the 100th condition 120, namely rain category 3 (lots of rain), fertilizer dose = 250 kg/ha N, 78 kg/ha P, and 156 kg/ha K. These two conditions are not only influenced by giving large doses of N fertilizer, but this can also be seen from the data to 83-98 in

category 2 and 3 rain conditions, fertilizer dose = 100 kg/ha N, 100 kg/ha P, and 200 kg/ha K. This data corresponds to the train value = 0.01 MSE (0.08 RMSE) and test value = 0.01 MSE (0.11 RMSE). The calculated RMSE value is low, indicating that the variation in values produced by a forecast model is close to the variation in the observed value. RMSE calculates how different a set of values are. The smaller the RMSE value, the closer the predicted and observed values are.

The MSE value obtained during the testing process demonstrates that the back propagation neural network can accurately forecast planting patterns [27]. The performance of these two variables can be improved by expanding the training data and adjusting network performance factors such as goal error, number of epochs, network design, type of activation function, and so on.

In the 30th-80th data, more fertilizer was given because farmers received subsidized fertilizer, so they were able to buy large quantities and at that time there was a lot of rain, so farmers were worried that a lot of fertilizer would be carried away by water flow due to rain, this is the basis for giving N fertilizer. more.

Applying 100, 150, 250, and 270 kg/Ha of fertilizer will produce optimal corn plant growth provided that there is sufficient water available in the form of monthly rainfall in the research location area (Nagrak District, Sukabumi Regency). As shown in table 4.5, the rain forecast for 2022 in table 5.

Table 5. Best Recommendations for Planting Corn

Month	Rainfall Category	Recommendation
January	High	Planting
February	High	Planting
March	Medium	Planting
April	Dry	Not Planting
May	Medium	Not Planting
June	Dry	Not Planting
July	Dry	Not Planting
August	Medium	Planting
September	High	Planting
October	High	Planting
November	High	Planting
December	High	Planting

It is hoped that these predictions can be used as a reference by corn farmers to start planting corn each planting season. The habit of farmers in Indonesia is to start planting corn in August – September, so that a big harvest in November is inevitable. With these forecast results, corn plants can be planted at the end of the rainy season, namely January - March and at the beginning of the rainy season, namely August – December.

4 Conclusion

After conducting research using the k-NN, Naïve Bayes, Decision Tree, and LSTM classification methods to estimate fertilizer application using a climate data dataset. The results of research using the k-NN method obtained 92% accuracy, using the Naive Bayes method obtained 100% accuracy, using the decision tree method obtained 88% accuracy, and using the LSTM method obtained 100% accuracy. From these results it can be concluded that the performance of the Naïve Bayes and LSTM methods is better than the k-NN and Decision Tree methods. This means that the application of N, P and K fertilizer to corn plants is influenced by the availability of rainwater. In rainfall categories 2 (moderate) and 3 (a lot), fertilizer can be absorbed by plants so that corn crop yields can be optimal, conversely if

sufficient water is not available then corn crop yields cannot be optimal. Predictive recommendations for corn planting in the next planting season is optimal planting at the end of the rainy season, namely January to March and at the beginning of the rainy season, namely August to December. The optimum fertilizer application is N fertilizer dose = 100-270 kg/ha, P fertilizer dose = 63-100 kg/ha, and K fertilizer dose = 156-200 kg/ha e.

References

1. M. Briga, E. Koetsier, J.J. Boonekamp, B. Jimeno, and S. Verhulst, Food availability affects adult survival trajectories depending on early developmental conditions in *Proceedings of the Royal Society B: Biological Sciences* **284**, 20162287 (2017)
2. B.N. Ekwueme and J.C. Agunwamba, Trend analysis and variability of air temperature and rainfall in regional river basins. *Civ. Eng. J.* **7**, 816 (2021)
3. L.S. Pereira, Water, agriculture and food: challenges and issues. *Water Resour. Manag.* **31**, 2985 (2017)
4. W. Jiao, L. Wang, W.K. Smith, Q. Chang, H. Wang, and P. D'Odorico, Observed increasing water constraint on vegetation growth over the last three decades. *Nat. Commun.* **12**, (2021)
5. C. Ingrao, R. Strippoli, G. Lagioia, and D. Huisinigh, Water scarcity in agriculture: An overview of causes, impacts and approaches for reducing the risks. *Heliyon* **9**, e18507 (2023)
6. Y. Apriyana *et al.*, The integrated cropping calendar information system: A coping mechanism to climate variability for sustainable agriculture in Indonesia. *Sustainability* **13**, 6495 (2021)
7. P.K. Dhillon and B. Tanwar, Rice bean: A healthy and cost-effective alternative for crop and food diversity. *Food Secur.* **10**, 525 (2018)
8. M.I. Alshelmani, E.A. Abdalla, U. Kaka, and M. Abdul Basit, Nontraditional feedstuffs as an alternative in poultry feed in *Advances in Poultry Nutrition Research* (IntechOpen, 2021)
9. H. Kebede, R. Sui, D.K. Fisher, K.N. Reddy, N. Bellaloui, and W.T. Molin, Corn yield response to reduced water use at different growth stages. *Agric. Sci.* **05**, 1305 (2014)
10. K.P. Panigrahi, H. Das, A.K. Sahoo, and S.C. Moharana, Maize leaf disease detection and classification using machine learning algorithms. in *Advances in intelligent systems and computing*, (Springer, 2020)
11. A.A. Qaffas, M.-A.B. HajKacem, C.-E.B. Ncir, and O. Nasraoui, An Explainable Artificial intelligence Approach for Multi-Criteria ABC item classification. *J. of Theoretical and Appl. Electron. Commer. Res.* **18**, 848 (2023)
12. M. Mahlayeye, R. Darvishzadeh, and A. Nelson, Cropping patterns of annual crops: A remote sensing review. *Remote. Sens.* **14**, 2404 (2022)
13. A. Nasrallah *et al.*, Sentinel-1 data for winter wheat phenology monitoring and mapping. *Remote. Sens.* **11**, 2228 (2019)
14. R. Valcarce-Diñeiro, B. Arias-Pérez, J.M. Lopez-Sanchez, and N. Sánchez, Multi-temporal dual- and quad-polarimetric synthetic aperture radar data for CROP-type mapping. *Remote. Sens.* **11**, 1518 (2019)
15. G. A. Abubakar *et al.*, Mapping maize fields by using Multi-Temporal Sentinel-1A and Sentinel-2A images in Makarfi, northern Nigeria, Africa. *Sustainability* **12**, 2539 (2020)

16. L. Gao, J. Song, X. Liu, J. Shao, J. Liu, and J. Shao, Learning in high-dimensional multimedia data: the state of the art. *Multimed. Syst.* **23**, 303 (2015)
17. F. Masood, W.U. Khan, S.U. Jan, and J. Ahmad, AI-Enabled traffic control prioritization in Software-Defined IoT networks for smart agriculture. *Sens.* **23**, no. 19, 8218 (2023)
18. M. Ozcan and S. Peker, A classification and regression tree algorithm for heart disease modeling and prediction. *Healthc. Analytics* **3**, 100130 (2022)
19. S. Kalogiannidis, Impact of employee motivation on organizational performance. a scoping review paper for public sector. *Strategic J. of Bus. & Ch. Manag.* **8**, (2021)
20. K.S.K. Patro, V.K. Yadav, V.S. Bharti, A. Sharma, A. Sharma, and T. Senthilkumar, IoT and ML approach for ornamental fish behaviour analysis. *Scientific Reports* **13**, (2023)
21. S. Chatfield, Recommendations for secondary analysis of qualitative data. *The Qualitative Rep.* **25**, 833 (2020)
22. D. Iskanto, Role of products element in determining decisions of purchase. *Inovbiz J. Inov. Bis.* **8**, 200 (2020)
23. K. Reichardt and L.C. Timm, How plants absorb nutrients from the soil. in *Springer eBooks*, (Springer, 2019)
24. R.T. Handayanto and H. Herlawati, Machine learning berbasis desktop dan web dengan metode jaringan syaraf tiruan untuk sistem pendukung keputusan. *J. Komtika (Komp. Dan Inform.)* **4**, 15 (2020)
25. A. Alwiyah, T.T. Louangdy, and A. Yolandari, Relation of relationship between research theory and variable with management case study. *Aptisi Transac. on Manag. (ATM)* **2**, 70 (2018)
26. G. Naidu, T. Zuva, and E.M. Sibanda, A review of evaluation metrics in Machine learning Algorithms. *Lect. Not. in Net. and Syst.* **724**, 15 (2023)
27. S.B. Sangeetha, N.R.W. Blessing, N. Yuvaraj, and J.A. Sneha, Improving the training pattern in Back-Propagation neural networks using Holt-Winters' seasonal method and gradient boosting model. in *Algorithms for intelligent systems*, (Springer, 2020)