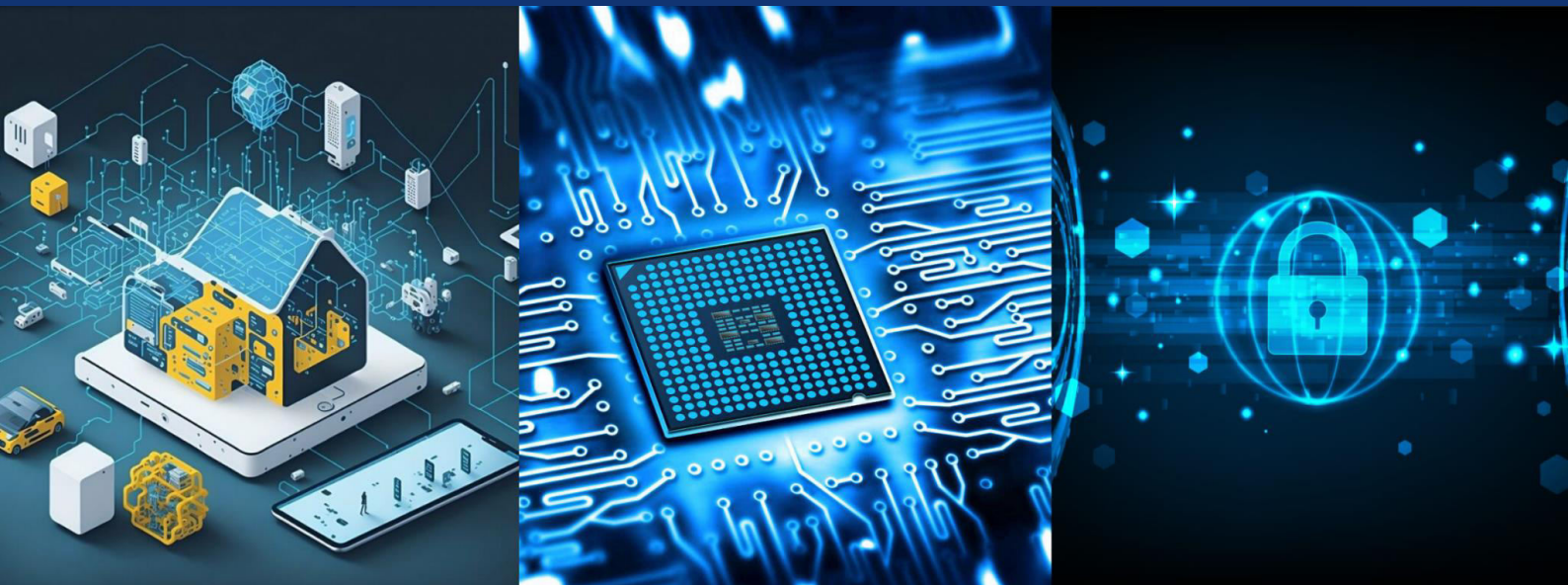
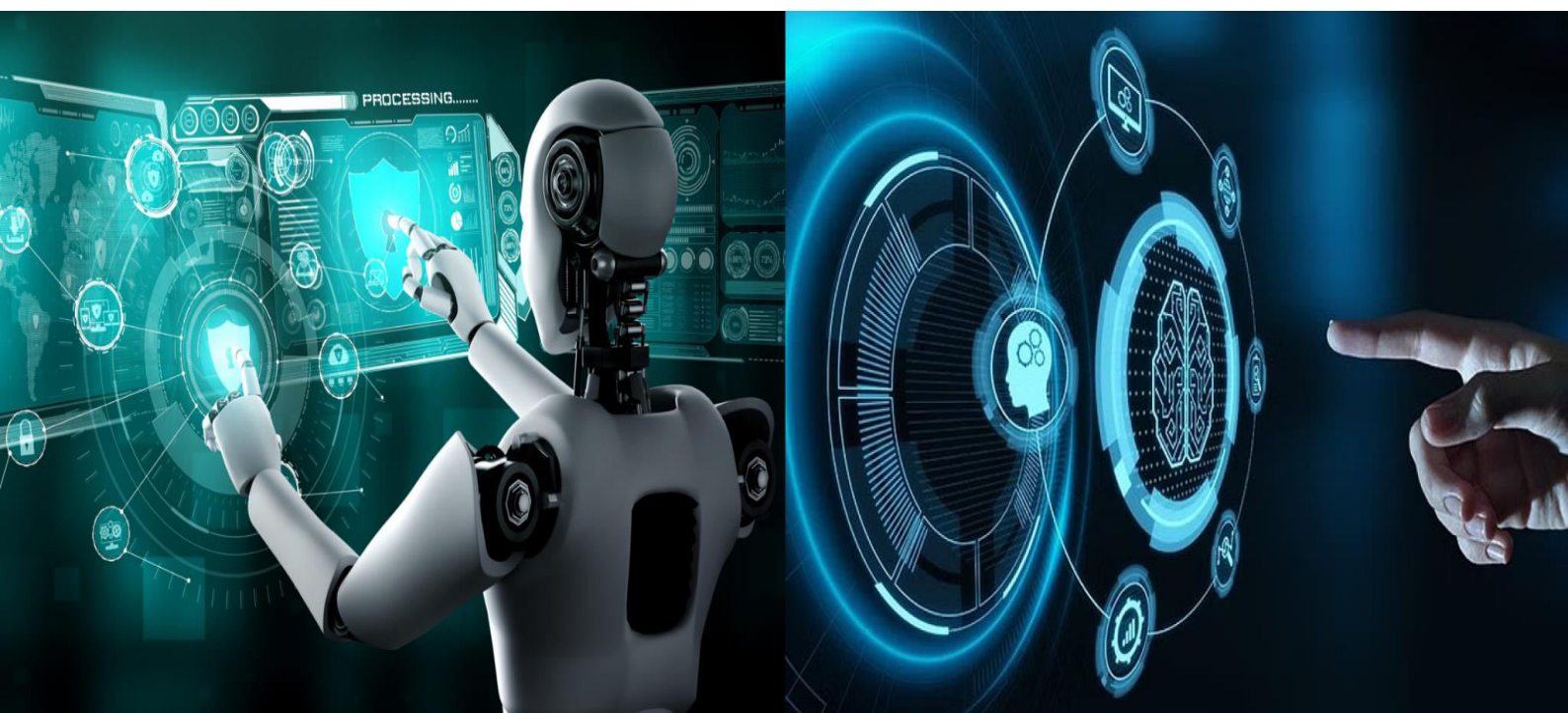




International Journal of Innovative Research in Computer and Communication Engineering

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)





International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

Efficient Crop Yield Prediction and Climate - Based Recommendation Using Machine Learning

S. Josphine¹, A. Nithin Venkata Naga Sai², P. Sowmya Naga Sri³, P. Dona Sunanda⁴,
U. Sai Gowtham⁵

Assistant Professor, Department of CSE, Eluru College of Engineering & Technology, Eluru, India¹

B. Tech Student, Department of CSE, Eluru College of Engineering & Technology, Eluru, India^{2,3,4,5}

ABSTRACT: The growth of agriculture is significantly influenced by various soil parameters, including Nitrogen, Phosphorus, Potassium, crop rotation, soil moisture, pH, surface temperature, and weather factors such as temperature and rainfall. The integration of technology in agriculture can enhance crop productivity, leading to improved yields for farmers. This paper aims to offer a solution for Smart Agriculture by monitoring agricultural fields, thereby assisting farmers in significantly boosting their productivity. We present a system in the form of a website that employs Machine Learning techniques to predict the most profitable crops based on current weather and soil conditions. Additionally, this system can forecast crop yields by analyzing weather parameters, soil characteristics, and historical crop yields. Ultimately, the project develops a comprehensive system that integrates data from multiple sources, data analytics, and predictive analysis to enhance crop yield productivity and increase farmers' profit margins over the longterm.

KEYWORDS: machine Learning, Prediction, Analyze and yield.

I. INTRODUCTION

Agriculture has an extensive history in India. Recently, India is ranked second in the farm output worldwide [15]. Agriculture-related industries such as forestry and fisheries contributed for 16.6% of 2009 GDP and around 50% of the total workforce. Agriculture's monetary contribution to India's GDP is decreasing [1]. The crop yield is the significant factor contributing in agricultural monetary. The crop yield depends on multiple factors such as climatic, geographic, organic, and financial elements [6]. It is difficult for farmers to decide when and which crops to plant because of fluctuating market prices [7]. Citing to Wikipedia figures India's suicide rate ranges from 1.4-1.8% per 100,000 populations, over the last 10 years [15]. Farmers are unaware of which crop to grow, and what is the right time and place to start due to uncertainty in climatic conditions. The usage of various fertilizers is also uncertain due to changes in seasonal climatic conditions and basic assets such as soil, water, and air. In this scenario, the crop yield rate is steadily declining [2]. The solution to the problem is to provide a smart userfriendly recommender system to the farmers.

The crop yield prediction is a significant problem in the agriculture sector [3]. Every farmer tries to know crop yield and whether it meets their expectations [4], thereby evaluating the previous experience of the farmer on the specific crop predict the yield [3]. Agriculture yields rely primarily on weather conditions, pests, and preparation of harvesting operations. Accurate information on crop history is critical for making decisions on agriculture risk management [5].

In this paper, we have proposed a model that addresses these issues. The novelty of the proposed system is to guide the farmers to maximize the crop yield as well as suggest the most profitable crop for the specific region. The proposed model provides crop selection based on economic and environmental conditions, and benefit to maximize the crop yield that will subsequently help to meet the increasing demand for the country's food supplies [8]. The proposed model predicts the crop yield by studying factors such as rainfall, temperature, area, season, soil type etc.

The system also helps to determine the best time to use fertilizers. The existing system which recommends crop yield is either hardware based being costly to maintain, or not easily accessible. The proposed system suggests a mobile-based application that precisely predicts the most profitable crop by predicting the crop yield. The use of GPS helps to identify the user location. The user provides an area under cultivation and soil type as inputs. According to the



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

requirement, the model predicts the crop yield for a specific crop. The model also recommends the most profitable crop and suggests the right time to use the fertilizers .

The major contributions of the paper are enlisted below

- Prediction of the crop yield for specific regions by executing various Machine Learning algorithms, with a comparison of error rate and accuracy.
- A user-friendly mobile application to recommend the most profitable crop.
- A GPS based location identifier to retrieve the rainfall estimation at the given area.
- A recommender system to suggest the right time for using fertilizers.

The organization of the rest of the paper is as follows: Section II discusses the background work of researchers in the field of agriculture and yield prediction. Section III presents the proposed model for yield prediction and recommends which crop for cultivation. The model also suggests the best suitable time for the use of fertilizers. Section IV discusses the results and Section V concludes the paper.

1.1 MOTIVATION

The history of agriculture in India dates back to the Indus Valley Civilization Era. India ranks second in this sector. Agriculture and allied sectors like forestry and fisheries account for 15.4 percent of the GDP (gross domestic product) with about 31 percent of the workforce. India ranks first globally with the highest net cropped area followed by US and China. Agriculture is demographically the broadest economic sector and plays a significant role in the overall socioeconomic fabric of India. Due to the revolution in industrialization, the economic contribution of agriculture to India's GDP is steadily declining with the country's broad-based economic growth.

1.2 PROBLEM DEFINITION

The problem that the Indian Agriculture sector is facing is the integration of technology to bring the desired outputs. With the advent of new technologies and overuse of non-renewable energy resources patterns of rainfall and temperature are disturbed. The inconsistent trends developed from the side effects of global warming make it cumbersome for the farmers to clearly predict the temperature and rainfall patterns thus affecting their crop yield productivity. In order to perform accurate prediction and handle inconsistent trends in temperature and rainfall various machine learning algorithms like Recurrent Neural Network (RNN), Long Short-Term Memory (LSTM) can be applied to get a pattern. In past, many researchers have applied machine learning techniques to enhance agricultural growth of the country.

1.3 OBJECTIVE OF THE PROJECT

This paper focuses on predicting the yield of the crop by applying various machine learning techniques. The outcome of these techniques is compared on the basis of mean absolute error. The prediction made by machine learning algorithms will help the farmers to decide which crop to grow to get the maximum yield by considering factors like temperature, rainfall, area, etc. The main objective of this project is to develop a system that predicts crop yields and recommends the most suitable fertilizers for a given crop. The system aims to assist farmers in making informed decisions about fertilizer application, ultimately leading to increased crop productivity and reduced costs.

II. LITERATURE SURVEY

Title: Predicting Yield of the Crop Using Machine Learning Algorithm:

The agriculture sector is crucial to a country's economy, but climate and environmental changes pose significant threats. Machine learning (ML) offers practical solutions. This paper focuses on predicting crop yield using historical data (weather, soil, and crop yield) and the Random Forest algorithm. Real data from Tamil Nadu was used to build and test the models, enabling farmers to predict crop yield before cultivation.

Title: Applications of Machine Learning Techniques in Agricultural Crop Production:

This paper reviews research on machine learning techniques in agricultural crop production. Accurate crop production forecasts are crucial for policy decisions, but current estimates are subjective. This study explores statistically sound, objective prediction methods using machine learning. Various techniques, including artificial neural networks, decision trees, and support vector machines, are applied to crop yield evaluation.



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

Title: A Model for Prediction of Crop Yield:

Yield prediction is crucial in agriculture. Data mining techniques can help predict crop yield using historical data. This research proposes a system using association rule mining to predict crop yield. The model is tested on data from Tamil Nadu, India, and the results show efficient prediction of crop yield production.

Title: Agricultural Crop Yield Prediction Using Artificial Neural Network Approach:

Weather conditions significantly impact crop yield. Artificial Neural Networks (ANNs) can model and predict crop yield by considering various parameters such as soil type, pH, temperature, rainfall, and humidity. This research uses ANNs to predict crop yield and identify suitable crops based on these parameters.

Title: Predictive Ability of Machine Learning Methods for Massive Crop Yield Prediction:

Accurate crop yield estimation is crucial for agricultural planning. This study compares the predictive accuracy of various machine learning (ML) techniques, including linear regression, for crop yield prediction in ten crop datasets. The results show that M5-Prime and k-nearest neighbor techniques achieve the lowest errors and highest correlation factors, making M5-Prime a suitable tool for massive crop yield prediction.

III. SYSTEM ANALYSIS

3.1 Existing System

The current agricultural system is plagued by numerous challenges, including soil erosion, climate change, and financial strains, which have made it increasingly difficult for farmers to sustain profitability and production. However, the integration of machine learning (ML) approaches has emerged as a viable solution, providing data-driven insights and suggestions to enhance agricultural sustainability and productivity. Various ML techniques, such as linear regression, Random Forest, Support Vector Regression (SVR), neural networks, and deep learning, have been successfully applied in crop recommendation systems, yield prediction, disease and pest detection, and precision farming. These models have demonstrated promising results, including improved crop yields, enhanced crop health, optimized agricultural inputs, and increased water usage efficiency. Nevertheless, challenges persist, including data variability, model interpretability, infrastructural issues, and ethical concerns, which must be addressed to fully harness the potential of ML in agriculture.

3.3.1 Disadvantages:

- Data Variability: ML models require high-quality and consistent data, which can be challenging to obtain in agriculture due to varying weather conditions, soil types, and crop varieties.
- Model Interpretability: ML models can be complex and difficult to interpret, making it challenging for farmers to understand the decisionmaking process.
- Infrastructural Issues: Many farmers lack access to the necessary infrastructure, such as high-speed internet, computers, and sensors, to implement ML solutions.
- Ethical Concerns: The use of ML in agriculture raises ethical concerns, such as the potential for biased decision-making, data privacy issues, and the impact on small-scale farmers.
- High Cost: Implementing ML solutions can be expensive, making it challenging for small-scale farmers to adopt these technologies.
- Limited Domain Expertise: ML models require domain expertise to develop and implement, which can be a challenge in agriculture where experts may not have the necessary technical skills.
- Dependence on Technology: ML solutions can be dependent on technology, which can be unreliable or unavailable in rural areas.

3.2 PROPOSED SYSTEM

Most of the Existing models utilized Neural networks, random forests, KNN regression techniques for CYP and a variety of ML techniques were also used for best prediction. The problems faced in existing research for crop yield prediction using machine learning are stated below:

- Creation, repair and maintenance of ML algorithms required huge costs as they are very complex.
- ML technique used for Crop yield prediction(mustard, wheat) combined input and output data but failed to obtain better results statistically



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

- Due to the nature of linear connection in the parameters, the regression model was failed to provide the exact prediction in a complex situation such as extreme value data and nonlinear data.
- The existing K-NN models were used for classification for yield prediction but lowered the
- performance due to nonlinear and highly adaptable issues present in KNN. They were operated in a locality model that incremented the dimensionality of the input vector made confusion for classification.
- An appropriate decision was not taken during classification because a fewer quantity of data was available for estimation of crop yield.

3.2.1 Advantages:

1. From the studies most of the common algorithms used were CNN, LSTM, DNN algorithms but still improvement was still required further in CYP.
2. The present research shows several existing models that consider elements such as temperature, weather condition, performing models for the effective crop yield prediction.

Ultimately, the experimental study showed the combination of ML with the agricultural domain field for improving the advancement in crop prediction.

3.3 MODULES

3.3.1 SERVICE PROVIDER

In this module, the Service Provider has to login by using valid user name and password. After login successful he can do some operations such as Browse Agriculture Data Sets and Train & Test, View Trained and Tested Accuracy in Bar Chart, View Trained and Tested Accuracy Results, View All Crop Yield and Production Prediction, View All Crop Recommendations, Download Predicted Data Sets, View All Remote Users, View Crop Yield Prediction Per Acre Results.

3.3.2 VIEW AND AUTHORIZE USERS

In this module, the admin can view the list of users who all registered. In this, the admin can view the user's details such as, user name, email, address and admin authorizes the users.

3.3.3 REMOTE USER

In this module, there are n numbers of users are present. User should register before doing any operations. Once user registers, their details will be stored to the database. After registration successful, he has to login by using authorized user name and password. Once Login is successful user will do some operations like predict crop yield and production, predict crop recommendation, view your profile.

IV. SYSTEM DESIGN

4.1 SYSTEM ARCHITECTURE

The system architecture for an E-Learning platform involves several interconnected layers that facilitate effective virtual learning and knowledge sharing. At the user interface layer, a responsive web portal and mobile apps provide seamless access to learning materials and features. The application layer incorporates modules for user authentication, content management, interaction, and personalization, utilizing advanced tools to tailor learning experiences to individual needs. The backend layer includes a database management system for storing user data and learning materials, as well as a file storage system for multimedia content. The cloud infrastructure ensures scalability, performance, and reliability through cloud hosting and content delivery networks. The security layer implements data encryption and access control to protect user information. Additionally, the analytics layer tracks user activity and monitors system performance to optimize the learning experience. Integration with third-party services adds extra functionalities like payment gateways and email services. This architecture ensures that the E-Learning platform is accessible, efficient, and engaging, providing a robust environment for learners to gain and share knowledge.



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

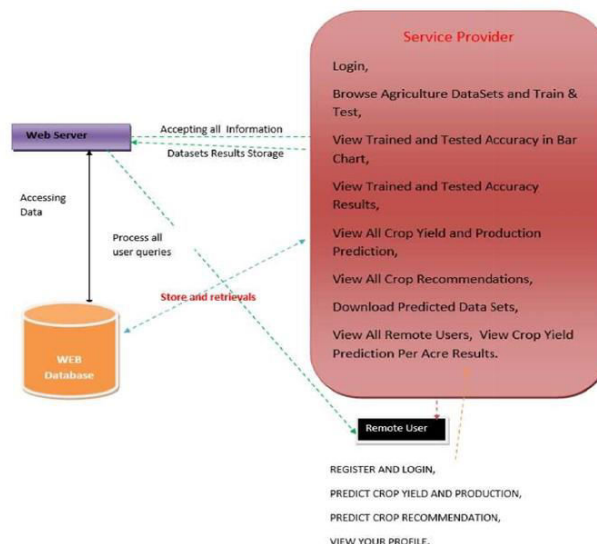


Fig 1: System Architecture

4.2 UML DIAGRAMS

UML stands for Unified Modeling Language. UML is a standardized general-purpose modeling language in the field of object-oriented software engineering. The standard is managed, and was created by, the Object Management Group. The goal is for UML to become a common language for creating models of object oriented computer software. In its current form UML is comprised of two major components: a Meta-model and a notation. In the future, some form of method or process may also be added to or associated with, UML. The Unified Modeling Language is a standard language for specifying, Visualization, Constructing and documenting the artifacts of software system, as well as for business modeling and other non-software systems.

UML was created as a result of the chaos revolving around software development and documentation. In the 1990s, there were several different ways to represent and document software systems.

The UML represents a collection of best engineering practices that have proven successful in the modeling of large and complex systems. The UML is a very important part of developing objects oriented software and the software development process. The UML uses mostly graphical notations to express the design of software projects.

4.2a GOALS:

The Primary goals in the design of the UML are as follows:

1. Provide users a ready-to-use, expressive visual modeling Language so that they can develop and exchange meaningful models.
2. Provide extendibility and specialization mechanisms to extend the core concepts.
3. Be independent of particular programming languages and development process.
4. Provide a formal basis for understanding the modeling language.
5. Encourage the growth of object oriented tools market.
6. Support higher level development concepts such as collaborations, frameworks, patterns and components.
7. Integrate best practices.

V. RESULTS

5.1 ALGORITHMS

5.1.1 DECISION TREE CLASSIFIERS:

Decision tree classifiers are used successfully in many diverse areas. Their most important feature is the capability of capturing descriptive decision making knowledge from the supplied data. Decision tree can be generated from training



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

sets. The procedure for such generation based on the set of objects (S), each belonging to one of the classes $C_1, C_2, \dots, C_k$ is as follows:

Step 1. If all the objects in S belong to the same class, for example C_i , the decision tree for S consists of a leaf labeled with this class

Step 2. Otherwise, let T be some test with possible outcomes $O_1, O_2, \dots, O_n$.

Each object in S has one outcome for T so the test partitions S into subsets $S_1, S_2, \dots, S_n$ where each object in S_i has outcome O_i for T. T becomes the root of the decision tree and for each outcome O_i we build a subsidiary decision tree by invoking the same procedure recursively on the set S_i .

5.1.2 GRADIENT BOOSTING

Gradient boosting is a machine learning technique used in regression and classification tasks, among others. It gives a prediction model in the form of an ensemble of weak prediction models, which are typically decision trees.[1][2] When a decision tree is the weak learner, the resulting algorithm is called gradient-boosted trees; it usually outperforms random forest. A gradient-boosted trees model is built in a stage-wise fashion as in other boosting methods, but it generalizes the other methods by allowing optimization of an arbitrary differentiable loss function.

5.1.3 K-NEAREST NEIGHBORS (KNN)

- Simple, but a very powerful classification algorithm
- Classifies based on a similarity measure
- Non-parametric
- Lazy learning
- Does not “learn” until the test example is given
- Whenever we have a new data to classify, we find its K-nearest neighbors from the training data.

Example:

- Training dataset consists of k-closest examples in feature space
- Feature space means, space with categorization variables (non-metric variables)
- Learning based on instances, and thus also works lazily because instance close to the input vector for test or prediction may take time to occur in the training dataset

5.1.4 LOGISTIC REGRESSION CLASSIFIERS

Logistic regression analysis studies the association between a categorical dependent variable and a set of independent (explanatory) variables. The name logistic regression is used when the dependent variable has only two values, such as 0 and 1 or Yes and No. The name multinomial logistic regression is usually reserved for the case when the dependent variable has three or more unique values, such as Married, Single, Divorced, or Widowed. Although the type of data used for the dependent variable is different from that of multiple regression, the practical use of the procedure is similar.

Logistic regression competes with discriminant analysis as a method for analyzing categorical-response variables. Many statisticians feel that logistic regression is more versatile and better suited for modeling most situations than is discriminant analysis. This is because logistic regression does not assume that the independent variables are normally distributed, as discriminant analysis does.

This program computes binary logistic regression and multinomial logistic regression on both numeric and categorical independent variables. It reports on the regression equation as well as the goodness of fit, odds ratios, confidence limits, likelihood, and deviance. It performs a comprehensive residual analysis including diagnostic residual reports and plots. It can perform an independent variable subset selection search, looking for the best regression model with the fewest independent variables. It provides confidence intervals on predicted values and provides ROC curves to help determine the best cutoff point for classification. It allows you to validate your results by automatically classifying rows that are not used during the analysis.

5.1.5 NAÏVE BAYES

The naive bayes approach is a supervised learning method which is based on a simplistic hypothesis: it assumes that the presence (or absence) of a particular feature of a class is unrelated to the presence (or absence) of any other feature .



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

Yet, despite this, it appears robust and efficient. Its performance is comparable to other supervised learning techniques. Various reasons have been advanced in the literature. In this tutorial, we highlight an explanation based on the representation bias. The naive bayes classifier is a linear classifier, as well as linear discriminant analysis, logistic regression or linear SVM (support vector machine). The difference lies on the method of estimating the parameters of the classifier (the learning bias).

While the Naive Bayes classifier is widely used in the research world, it is not widespread among practitioners which want to obtain usable results. On the one hand, the researchers found especially it is very easy to program and implement it, its parameters are easy to estimate, learning is very fast even on very large databases, its accuracy is reasonably good in comparison to the other approaches. On the other hand, the final users do not obtain a model easy to interpret and deploy, they does not understand the interest of such a technique.

Thus, we introduce in a new presentation of the results of the learning process. The classifier is easier to understand, and its deployment is also made easier. In the first part of this tutorial, we present some theoretical aspects of the naive bayes classifier. Then, we implement the approach on a dataset with Tanagra. We compare the obtained results (the parameters of the model) to those obtained with other linear approaches such as the logistic regression, the linear discriminant analysis and the linear SVM. We note that the results are highly consistent. This largely explains the good performance of the method in comparison to others. In the second part, we use various tools on the same dataset (Weka 3.6.0, R 2.9.2, Knime 2.1.1, Orange 2.0b and RapidMiner 4.6.0).

5.1.6 RANDOM FOREST

Random forests or random decision forests are an ensemble learning method for classification, regression and other tasks that operates by constructing a multitude of decision trees at training time. For classification tasks, the output of the random forest is the class selected by most trees. For regression tasks, the mean or average prediction of the individual trees is returned. Random decision forests correct for decision trees' habit of overfitting to their training set. Random forests generally outperform decision trees, but their accuracy is lower than gradient boosted trees. However, data characteristics can affect their performance.

The first algorithm for random decision forests was created in 1995 by Tin Kam Ho[1] using the random subspace method, which, in Ho's formulation, is a way to implement the "stochastic discrimination" approach to classification proposed by Eugene Kleinberg.

An extension of the algorithm was developed by Leo Breiman and Adele Cutler, who registered "Random Forests" as a trademark in 2006 (as of 2019, owned by Minitab, Inc.). The extension combines Breiman's "bagging" idea and random selection of features, introduced first by Ho[1] and later independently by Amit and Geman[13] in order to construct a collection of decision trees with controlled variance.

Random forests are frequently used as "blackbox" models in businesses, as they generate reasonable predictions across a wide range of data while requiring little configuration.

5.1.7 SVM

In classification tasks a discriminant machine learning technique aims at finding, based on an independent and identically distributed (iid) training dataset, a discriminant function that can correctly predict labels for newly acquired instances. Unlike generative machine learning approaches, which require computations of conditional probability distributions, a discriminant classification function takes a data point x and assigns it to one of the different classes that are a part of the classification task. Less powerful than generative approaches, which are mostly used when prediction involves outlier detection, discriminant approaches require fewer computational resources and less training data, especially for a multidimensional feature space and when only posterior probabilities are needed. From a geometric perspective, learning a classifier is equivalent to finding the equation for a multidimensional surface that best separates the different classes in the feature space.

SVM is a discriminant technique, and, because it solves the convex optimization problem analytically, it always returns the same optimal hyperplane parameter—in contrast to genetic algorithms (GAs) or perceptrons, both of which are widely used for classification in machine learning. For perceptrons, solutions are highly dependent on the initialization and termination criteria. For a specific kernel that transforms the data from the input space to the feature space, training



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

returns uniquely defined SVM model parameters for a given training set, whereas the perceptron and GA classifier models are different each time training is initialized. The aim of GAs and perceptrons is only to minimize error during training, which will translate into several hyperplanes' meeting this requirement.

The following figures present the sequence of screenshots of the results.



Fig 2a: Login Page

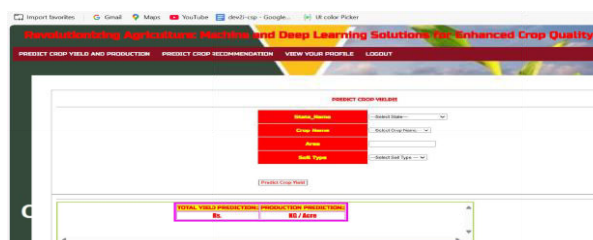


Fig 2b: User Details

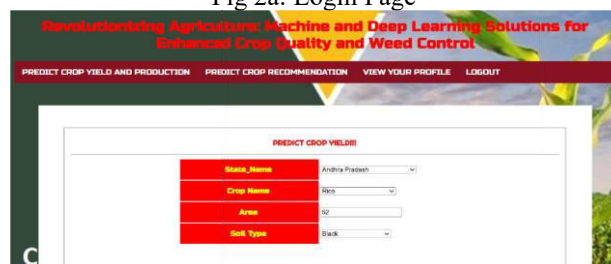


Fig 2c: Input the Attributes for yield prediction

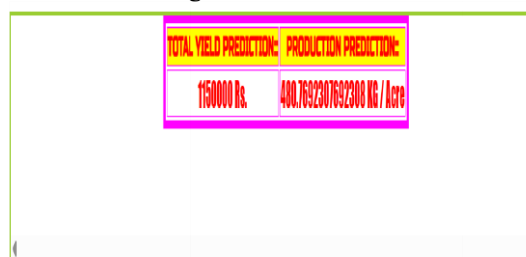


Fig 2d: Crop Yield Prediction.

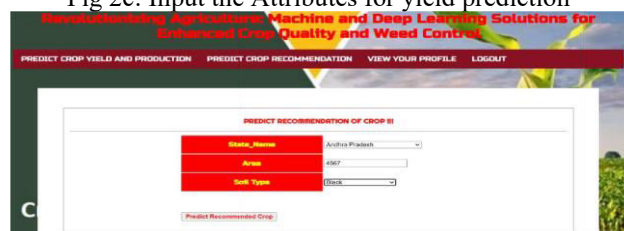


Fig 2e: Input the Attributes for Crop Recommendation



Fig 2f: Crop Recommendation

VI. CONCLUSIONS AND FUTURE WORK

6.1 CONCLUSIONS

This paper highlighted the limitations of current systems and their practical usage in yield prediction. It then introduced a viable yield prediction system designed for farmers, which offers connectivity via a mobile application. The mobile application includes multiple features that assist farmers in crop selection. The inbuilt predictor system helps farmers estimate the yield of a given crop, while the recommender system enables users to explore potential crops and their expected yields, allowing for more informed decision-making.

To achieve high accuracy in yield prediction, various machine learning algorithms, including Random Forest, Artificial Neural Networks (ANN), Support Vector Machines (SVM), Multiple Linear Regression (MLR), and K-Nearest Neighbors (KNN), were implemented and tested on datasets from Maharashtra and Karnataka. Comparative analysis of these algorithms indicated that Random Forest Regression outperformed the others, achieving an accuracy of 95%. Additionally, the proposed model explored the timing of fertilizer application and recommended appropriate durations.

6.2 FUTURE WORK

Future enhancements will focus on updating datasets periodically to ensure more accurate predictions and automating various processes within the system. One key improvement will be the implementation of a feature that recommends the correct type of fertilizer based on the crop and location. To achieve this, a thorough study of available fertilizers and



International Journal of Innovative Research in Computer and Communication Engineering (IJIRCCE)

(A Monthly, Peer Reviewed, Refereed, Scholarly Indexed, Open Access Journal)

their relationship with soil and climate conditions will be conducted. Furthermore, an in-depth analysis of statistical data will be performed to refine and enhance the predictive capabilities of the system.

REFERENCES

1. Umamaheswari S, Sreeram S, Kritika N, Prasanth DJ, "BioT: Blockchain-based IoT for Agriculture", 11th International Conference on Advanced Computing (ICoAC), 2019 Dec 18 (pp. 324-327). IEEE.
2. Jain A. "Analysis of growth and instability in the area, production, yield, and price of rice in India", Journal of Social Change and Development, 2018;2:46-66
3. Manjula E, Djodiltachoumy S, "A model for prediction of crop yield" International Journal of Computational Intelligence and Informatics, 2017 Mar;6(4):2349-6363.
4. Sagar BM, Cauvery NK., "Agriculture Data Analytics in Crop Yield Estimation: A Critical Review", Indonesian Journal of Electrical Engineering and Computer Science, 2018 Dec;12(3):1087-93.
5. Wolfert S, Ge L, Verdouw C, Bogaardt MJ, "Big data in smart farming– a review. Agricultural Systems", 2017 May 1;153:69-80.
6. Jones JW, Antle JM, Basso B, Boote KJ, Conant RT, Foster I, Godfray HC, Herrero M, Howitt RE, Janssen S, Keating BA, "Toward a new generation of agricultural system data, models, and knowledge products: State of agricultural systems science. Agricultural systems", 2017 Jul 1;155:269-88.
7. Johnson LK, Bloom JD, Dunning RD, Gunter CC, Boyette MD, Creamer NG, "Farmer harvest decisions and vegetable loss in primary production. Agricultural Systems", 2019 Nov 1;176:102672.
8. Kumar R, Singh MP, Kumar P, Singh JP, "Crop Selection Method to maximize crop yield rate using a machine learning technique", International conference on smart technologies and management for computing, communication, controls, energy, and materials (ICSTM), 2015 May 6 (pp. 138-145). IEEE.
9. Sriram Rakshith.K, Dr.Deepak.G, Rajesh M, Sudharshan K S, Vasanth S, Harish Kumar N, "A Survey on Crop Prediction using Machine Learning Approach", In International Journal for Research in Applied Science & Engineering Technology (IJRASET), April 2019, pp 3231- 3234.
10. Veenadhari S, Misra B, Singh CD, "Machine learning approach for forecasting crop yield based on climatic parameters", In 2014 International Conference on Computer Communication and Informatics, 2014 Jan 3 (pp. 1-5). IEEE.
11. Ghadge R, Kulkarni J, More P, Nene S, Priya RL, "Prediction of crop yield using machine learning", Int. Res. J. Eng. Technol. (IRJET), 2018.
12. Priya P, Muthaiah U, Balamurugan M, "Predicting yield of the crop using machine learning algorithm", International Journal of Engineering Sciences & Research Technology, 2018 Apr;7(1):1-7.
13. S. Pavani, Augusta Sophy Beulet P., "Heuristic Prediction of Crop Yield Using Machine Learning Technique", International Journal of Engineering and Advanced Technology (IJEAT), December 2019, pp 135-138.
14. <https://web.dev/progressive-web-apps/>
15. <https://www.wikipedia.org/>
16. Plewis I, "Analyzing Indian farmer suicide rates", Proceedings of the National Academy of Sciences, 2018 Jan 9;115(2): E117.
17. Nishant, Potnuru Sai, Pinapa Sai Venkat, Bollu Lakshmi Avinash, and B. Jabber. "Crop Yield Prediction based on Indian Agriculture using Machine Learning." In *2020 International Conference for Emerging Technology (INCET)*, pp. 1-4. IEEE, 2020.
18. Kale, Shivani S., and Preeti S. Patil. "A Machine Learning Approach to Predict Crop Yield and Success Rate." In *2019 IEEE Pune Section International Conference (PuneCon)*, pp. 1-5. IEEE, 2019.
19. Kumar, Y. Jeevan Nagendra, V. Spandana, V. S. Vaishnavi, K. Neha, and V. G. R. R. Devi. "Supervised Machine learning Approach for Crop Yield Prediction in Agriculture Sector." In *2020 5th International Conference on Communication and Electronics Systems (ICCES)*, pp.736-741. IEEE, 2020.
20. Nigam, Aruvansh, Saksham Garg, Archit Agrawal, and Parul Agrawal. "Crop yield prediction using machine learning algorithms." In *2019. Fifth International Conference on Image Information Processing (ICIIP)*, pp. 125-130. IEEE, 2019.
21. Bang, Shivam, Rajat Bishnoi, Ankit Singh Chauhan, Akshay Kumar Dixit, and Indu Chawla. "Fuzzy logic based crop yield prediction using temperature and rainfall parameters predicted through ARMA, SARIMA, and ARMAX models." In *2019 Twelfth International Conference on Contemporary Computing (IC3)*, pp. 1-6. IEEE, 2019.



INTERNATIONAL
STANDARD
SERIAL
NUMBER
INDIA



INTERNATIONAL JOURNAL OF INNOVATIVE RESEARCH

IN COMPUTER & COMMUNICATION ENGINEERING

 9940 572 462  6381 907 438  ijircce@gmail.com



www.ijircce.com

Scan to save the contact details