

Dual Stacking for Rapid Crop Yield Prediction Using LSTM-Based Analytical Deep Learning

Surbhi Tandon¹, Dr. Navneet Kaur²

¹Research Scholar, Computer Science & Applications, Desh Bhagat University. Email: surbhikh130@gmail.com

²Assistant Professor, Computer Science & Engg, Desh Bhagat University

Abstract: This paper presents a novel approach to rapid crop yield prediction through the integration of Dual Stacking and Long Short-Term Memory (LSTM)-based Analytical Deep Learning (ADL). The proposed DSA-LSTM framework leverages both short-term and long-term memory variables to efficiently forecast agricultural output. By utilizing Dual Stacking as a temporary repository during the Set Covering Approximation, the model enhances the accuracy and speed of yield prediction, addressing the complexities of agricultural data. The research explores various aspects of crop productivity, including environmental and soil factors, and applies advanced data mining techniques using WEKA to optimize attribute selection and clustering. Experimental results from the Hissar district in Haryana show the potential of this method for accurate yield forecasting. The study also emphasizes the importance of proper data mapping and fact-table management in enhancing prediction reliability. This methodology provides a scalable and adaptable solution for real-time agricultural decision-making, supporting stakeholders in improving crop management and productivity. The future scope of the research includes incorporating additional environmental, economic, and market factors for more precise predictions.

Keywords: Crop Yield Prediction, Dual Stacking, LSTM, Analytical Deep Learning, Set Covering Approximation, Data Mining, WEKA, Agricultural Forecasting

1. INTRODUCTION

Making predictions about crop yields is essential to ensuring that there is food security across the world, to providing help for agricultural planning, and to making it easier to foresee market conditions. When it comes to estimating crop output, traditional approaches typically struggle with speed and accuracy, particularly when it comes to combining varied agricultural data, such as soil health, meteorological conditions, crop variety, and cultivation procedures [1]. The accuracy and speed of crop production projections have both been significantly improved as a result of recent breakthroughs in deep learning methods, especially Long Short-Term Memory networks (LSTM networks). This work presents a novel Dual Stacking approach that is used with Analytical Deep Learning (ADL), which is based on Long Short-Term Memory (LSTM), in order to estimate agricultural yields more quickly and efficiently [2]. The Dual Stacking approach functions as a temporary repository, which makes the attribute selection process more efficient and improves the model's capacity to analyze large-scale agricultural data in an effective manner [3]. The technique efficiently utilizes both short-term and long-term memory factors, which are critical for properly projecting agricultural output by incorporating both present circumstances and previous patterns [4].

The Set Covering Approximation approach is examined in this study in order to optimize vertices that intersect with one another. This allows for the selection of the smallest possible set of characteristics, which is required for making yield forecasts that are correct. The practicality of this technique in real-world agricultural contexts is shown by its application to crop production forecasting in the Hissar area of Haryana [5]. This method makes the model more capable of making accurate predictions. This is accomplished by adding sophisticated data mining methods such as classification and clustering with WEKA, which results in optimized data structures and exact mapping [6]. This guarantees that the system will continue to be scalable and that it will be able to adapt to a variety of agricultural



settings. In addition to improving forecast speed and accuracy, the framework assists players in the agricultural sector in making well-informed choices based on data, which contributes to more environmentally friendly farming methods. This study lays the framework for future work, which may involve the integration of economic, market, and extra environmental aspects [7]. The development of a strong system that can be used worldwide to handle the difficulties that are presented by contemporary agriculture is the ultimate objective.

2. SET COVERING APPROXIMATION TO OPTIMIZE THE INTERSECTING VERTICES

During the process of ADL for issue modelling, it is essential to choose certain optimal variables based on the required interpretation, ensuring that the results of the user's query are delivered with maximum effectiveness and efficiency [6]. Given the multitude of variables present in agricultural data sets, it becomes essential to estimate and refine them effectively. To adhere to a specified greedy methodology recognised as the set covering approximation for enhancing and choosing overlapping vertices [15]. A significant category of optimisation challenges pertains to encompassing the minimal set of attributes necessary to forecast a concern associated with the user inquiry. The set cover approximation is a solution that falls within the realm of NP-completeness, and it can be elucidated in the following manner [7]:

“Set-Cover: Consider a domain X comprising n distinct points, alongside m subsets $S_1, S_2, \dots, S_m$ that include these points. The objective is to determine the minimum number of subsets required to encompass all the points. Additionally, the decision problem introduces a number k , enquiring whether it is feasible to cover the entirety of the points utilising k or fewer subsets.”

In the realm of set covering optimisation, we define a pair $\Sigma(X, S)$ as a vertex system, where $X = \{x_1, x_2, \dots, x_m\}$ represents a finite collection of attributes, and $S = \{s_1, s_2, \dots, s_n\}$ denotes a compilation of subsets derived from X , ensuring that each element within X is included in at least one subset from S . determine a collection $C = \{1, \dots, n\}$ with the least possible size that satisfies the following conditions [17]:

$$X = \bigcup_{i \in C} s_i$$

Approximation of set cover

In relation to our data collection, we will examine various sets pertinent to the agricultural sector along with their characteristics as detailed below:

$S_1 = \{\text{SoilTypes}, \text{pH}\},$

$S_2 = \{\text{ClimaticZone}, \text{Annual Rainfall (mm)}, \text{Wet period (months)}, \text{Vegetation}\}$

$S_3 = \{\text{Crop Name}, \text{Temperature}, \text{Rain-fall}, \text{Soil-Type}, \text{Producer}\}$

$S_4 = \{\text{Nutrient}, \text{Soil}, \text{Crop}, \text{Critical Level}, \text{Method}\}$

$S_5 = \{\text{Crop Name}, \text{Temperature}, \text{Rain-fall}\}$

$S_6 = \{\text{StateName}, \text{District_Name}, \text{Crop_Year}, \text{Season}, \text{Crop}, \text{Area}, \text{Production}\}$

$S_7 = \{\text{Crop Name, family class, Usable part, Duration (Days), industrial use, Domestic Use, Seed Use, Area cover worldwide, Production worldwide}\}$

$S_8 = \{\text{Year, District, Crop, Area, Tanks, Bore Wells, Open Wells, Production, Yield}\}$

$S_9 = \{\text{CROP-ZONES, AGRO-CLIMATIC REGIONS, SOIL TYPE, RAINFALL (Range in mm.)}\}$

$S_{10} = \{\text{District, Crop, Soil, N, P, K, Rain fall Range}\}$

$S_{11} = \{\text{SUBDIVISION, YEAR, Parameter, Jan, Feb, Mar, Apr, May, Jun, Jul, Aug, Sep, Oct, Nov, Dec, ANNUAL, JF, MAM, JJAS, OND}\}$

$S_{12} = \{\text{state, district, market, commodity, variety, arrival_date, min_price, max_price, modal_price}\}$

After approximation, a collection of optimized attributes required for HARYANA crop yield prediction as shown here:

Final Set = {District_Name, Crop_Year, Season, Crop, Area (Hect), Production (Tones), Yield (Tones/Hect), Min-Temperature, Max-Temperature, RH-Relative Humidity, Soil Type, Nitrogen -N, Phosphorous -P, Potassium -K, pH, Rain fall Range, CEC (Meq/100g), Duration in days, Wind speed km/h, productivity kg/hect}

3. META DATA CLASSIFICATION USING DUAL STAKING [OUR APPROACH]

In this endeavor, we are aiming to introduce an innovative strategy by employing the dual stacking technique to effectively map information through parallel processing during the prediction phase [12]. This document encompasses three primary types of dataset processing: the selection of fact tables, the implementation of dual stacking through various data mining methodologies, and the mapping of all attributes present in the associated fact tables. The entire sequence of the implementation process and its corresponding workflow is illustrated in figure [1]. It should possess both effectiveness and practicality when handling vast quantities of data, as well as the extent to which its influence on issue forecasting will ultimately be revealed [8].

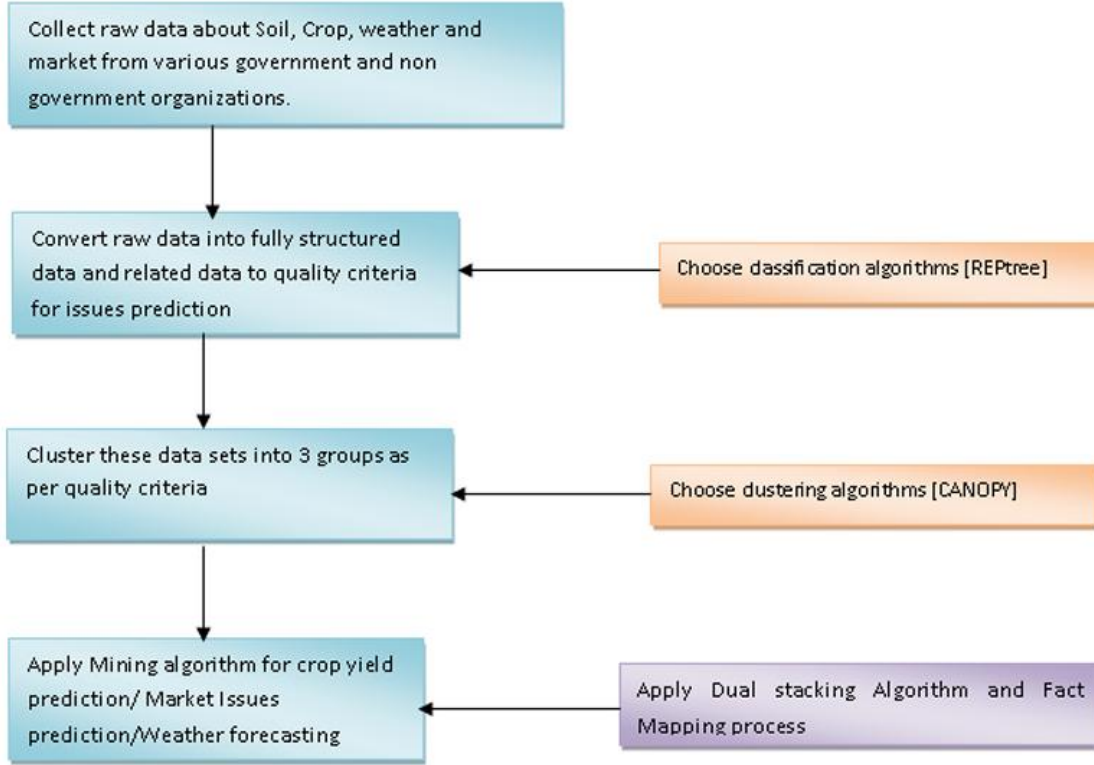


Figure 1: Block Architecture for implementation of Dual-Stack Algorithm

The illustration above features blue boxes that represent the foundational aspects of data model development, while the orange boxes depict the established algorithms employed for data analysis, classification, and clustering [13]. The presentation of the purple-colored box illustrates the connection to our innovative methodology [9].

4. ALGORITHMIC PROCEDURE OF IMPLEMENTATION

4.1. Integrated Approach

To ensure the effective execution of the dual-stack modelling approach for issue forecasting, which aims to facilitate decision-making for farmers, it is essential to comprehend three key algorithmic processes outlined here:

Dual Stacking: This technique involves the organisation of transient characteristics and their corresponding values related to the issues prediction model, allowing for the retrieval of properties and their values to facilitate simultaneous processing [15]. This process will set up two open-ended stacks [fig. 2] simultaneously to hold parameters for issuing and removing following knowledge prediction. The operation of this dual stack resembles that of the stack implementation [10]. Additionally, there exist two types of operations, namely pop and push, which facilitate the storage and retrieval of elements.

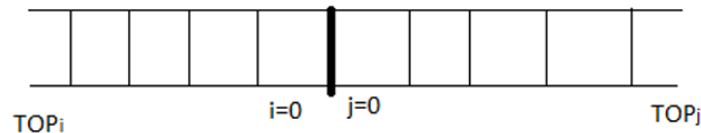


Figure 2: Dual Stack Implementation

Knowledge Support: This involves suggesting various parameters along with their associated values through a recommendation system tailored for farmers or agriculturalists. Additionally, it incorporates input derived from

existing knowledge to facilitate the mapping process with current fact tables [16]. This serves as a method for recommending and selecting options, enabling time efficiency and enhancing the optimisation of fact mapping [17].

Fact Mapping: The precision of the prediction model is significantly influenced by the effectiveness of the fact mapping process. Inaccurate data within the fact tables will lead to unreliable outcomes [11]. Consequently, the efficacy of this entire approach is significantly reliant on fact tables [18]. The aggregation of fact tables serves as our primary and crucial standard to ensure that efficiency is upheld.

The procedural stages for the cohesive methodology are delineated as follows:

- Retrieve characteristics from data tables associated with the problem for forecasting.
- Incorporate these characteristics along with their associated values into the dual-stack.
- Enhance the process by incrementing the counters i and j .
- Continue executing steps 1 through 3 until every associated data table has been fully retrieved.
- Remove all layered components and associate them with fact tables.
- Reduces the process counters i and j .
- Continue to execute steps 5 and 6 repeatedly until the mapping procedure reaches its conclusion.
- Anticipate the assertions and formulate the data structure in accordance with geographical coordinates of latitude and longitude.

Algorithmic steps for Dual Stacking

Each time an attempt is made to execute the strategy concerning crops, weather, and market dynamics, and when addressing challenging circumstances, adhere to these steps that are likely appropriate for dual stacking, as detailed during the selection and forecasting process:

1. Initialize common bottom=0
2. Initialize $TOP_i=0$
3. Initialize $TOP_j=0$
4. Choose features as per selection [crop, weather, market] for stacking.
5. Enter attribute_name into the left stack and its corresponding value in the right stack.
6. Increase TOP_i and TOP_j
7. Repeat step 4-6 till all the attributes are selected

At the time of selection:

At the time of prediction:

1. Find the current value of TOPi or TOPj

2. For count=TOPi to 0

Print dualstackname_left[count]

Print dualstackname_right[count]

Call mapping procedure with fact table for the attribute
dualstackname_left[count]

Return result

Decrease TOPi and TOPj

3. When the TOPi or TOPj became zero

4. Predict the report as per data variations.

In the context of multidimensional data modelling, this investigation gathers the subsequent insights derived from earlier research conducted by scholars in the field of agriculture. This research endeavour, akin to those iconic masterpieces, is succinctly elucidated here and will be elaborated upon in the forthcoming efforts.

5. CROP RELATED PREDICTIVE DATA MINING AND ANALYSIS USING WEKA

Prior to embarking on our analytical methodology for forecasting agricultural output, we will employ the current WEKA data mining algorithms and associated rules, thereby facilitating the presentation of a comparative analysis [12, 19]. Given the extensive size of the data set encompassing all crop-related records from Haryana, for the purpose of testing and experimentation, it is advisable to segment the data into district-specific .csv files. Begin the experiments with the initial data set associated with District Hissar.

5.1 STATISTICAL MEASUREMENT OF ATTRIBUTES

In any analytical examination, statistical metrics and their corresponding values serve as the primary benchmarks for discerning the identification of interrelations. In order to enhance the predictive data mining framework, it is essential to discern the connections between the various attributes within the agricultural sector. The table [1] presents information comprising 22 attributes and a total of 19,091 instances.

Table 1: Statistical measurement of attributes

Attribute	Type	Count	Distinct	Properties
District_name	Nominal	Missing 0%	51	
Crop_year	Numeric		13	Min 2005 Max 2018 X 2011 σ 687
Season	Nominal	Missing 0%	3	Kharif-151 Rabi -71 Whole Year 107
Crop	Nominal	28	42	

Area	Numeric	Unique 252 (77%)	8020	Min 1 Max 163378 X 19304.532 σ 37819.395
Production	Numeric	Missing 10 (3%)	8258	Min 0.01 Max 479988 X 19460.621 σ 53565.785
Yield	Numeric	Unique 294 (89%)	14876	Min 2005 Max 2018 X 2011 σ 2.15
Min Temperature	Numeric	Missing 0 (0%)	26	Min 9 Max 24.3 X 19.274 σ 5.874
Max Temperature	Numeric	Missing 0 (0%)	40	Min 32 Max 39.2 X 35.37 σ 534
Rh	Numeric	Missing 0 (0%)	41	Min 29 Max 69 X 48.04 σ 17.798
Soil_Type	Nominal	Missing 0 (0%)	5	
Nitrogen-N	Nominal	Missing 0 (0%)	2	
Phosphorous-P	Nominal	Missing 0 (0%)	11	
Potassium-K	Nominal	Missing 0 (0%)	8	
pH	Numeric	Missing 0 (0%)	5	
Rain Fall Range	Nominal	Missing 0 (0%)	9	
CEC	Nominal	Missing 0 (0%)	10	
Duration	Nominal	Missing 0 (0%)	7	
Wind Speed	Numeric	Missing 0 (0%)	1	

The illustration in figure [3] presents the statistical data pertaining to crop-related records, indicating that there are no gaps in the information. A total of 42 unique records of various crops are depicted across several districts in Haryana. In a similar vein, seek out the statistical data pertaining to additional characteristics such as seasonality, precipitation, soil composition, and so forth.

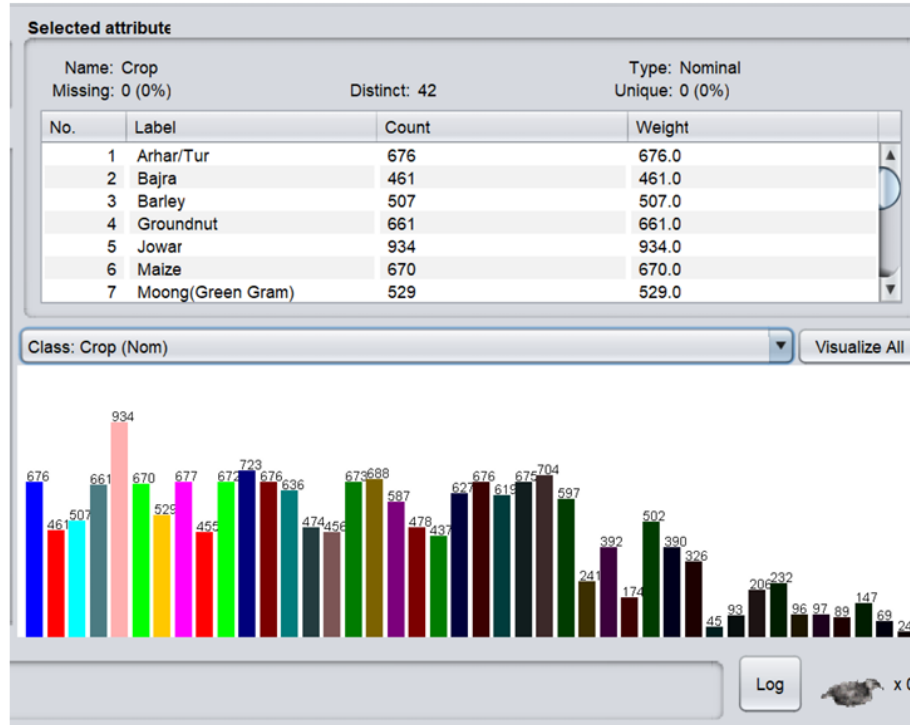


Figure 3: Screen shot of WEKA analysis”

5.2 WEKA Based Classifiers for Crop Issue

5.2.1 WEKA-Based Cluster Analysis Technique

Cluster Analysis serves as a technique in the realm of data mining, primarily utilised to reveal the arrangement of data and discern patterns within the core dataset. Data clustering denotes a systematic approach that categorises input data collections into separate groups, typically known as clusters. A cluster consists of elements that possess similar traits, indicating that every component within the cluster demonstrates a significant level of resemblance [21]. Clustering embodies a challenge in the domain of unsupervised learning, characterised by the absence of any prior knowledge apart from the input data sets at hand. The clustering sub-process has been employed in a diverse array of fields, including but not limited to image processing, video analysis, machine learning, image mining, biochemical assessments, bioinformatics applications, and various aspects of the petrochemical industry [22]. The result hinges on the attributes of the data collections and the intent behind the operation. The main aim of grouping dimensions is employed throughout diverse interrelated data collections. A wide array of techniques can be utilised for various clustering functions, leveraging hierarchical models, partitioned setups, graph-based frameworks, networks, trees, and numerous other types of aggregation and generalisation designed for particular relevant fields.

5.3 Considered Characteristics of Clustering After Analysis

Following an analysis grounded in WEKA, it becomes evident that certain essential factors regarding the data set must be taken into account to achieve precise clustering. Establishing significance criteria for the predictive parameters is crucial. The primary attributes to take into account regarding dual staking are outlined as follows:

Scalability- The intricacies of time and space must not overwhelm the extensive data sets we are examining for our cluster development utilising DSA.

Robustness- The DSA procedure ought to identify anomalies within the dataset, necessitating the reduction of attributes to enhance both speed and precision.

Order insensitivity- The DSA should not be overly sensitive to the point of altering the sequence of the data inputs.

Reduce the number of parameters defined by the user to ensure that data interpretation aligns more closely with the true value.

Data of various types can be expressed as numerical values, strings, or a combination of both, all managed through Data Structures and Algorithms.

Clusters of Irregular Form - The formations can take on various shapes, allowing for straightforward scaling in the future.

The principle of point proportion admissibility dictates that the act of duplicating a data set and subsequently re-clustering should not result in any alterations to the clustering outcomes.

A significant portion of clustering tasks necessitates repetitive processes to identify either locally or globally optimal solutions derived from high-dimensional datasets within the agricultural sector. Furthermore, actual data is depicted through a distinctive clustering approach, making the interpretations of these representations quite challenging. Consequently, it necessitates considerable effort to explore various algorithms or to manipulate distinct attributes within the same data set. As a result, these algorithms demand considerable computational resources, making the reduction of computational complexity a crucial concern for clustering techniques. Consequently, the parallel execution of clustering algorithms represents a highly effective strategy, and the parallel techniques will vary across different algorithms.

5.4 Our Work for Clustering Analytical Deep

Learning

Create a cluster within the design district that comprises a collection of analogous crop records in a distinctive manner.

The designated cluster for experimentation focusses on agricultural yield, emphasising the significance of attribute selection that exerts a greater influence.

Clusters exhibit a strong correlation concerning both the type of crops and the nature of the soil.

6. DUAL STACKING FOR ATTRIBUTION OF INTER-RELATION IDENTIFICATION

The execution of dual staking can be integrated within the WEKA mining tool, similar to the Bagging and Boosting functionalities offered by WEKA.

6.1 Meta Data Classification Using Staking [WEKA Defined]

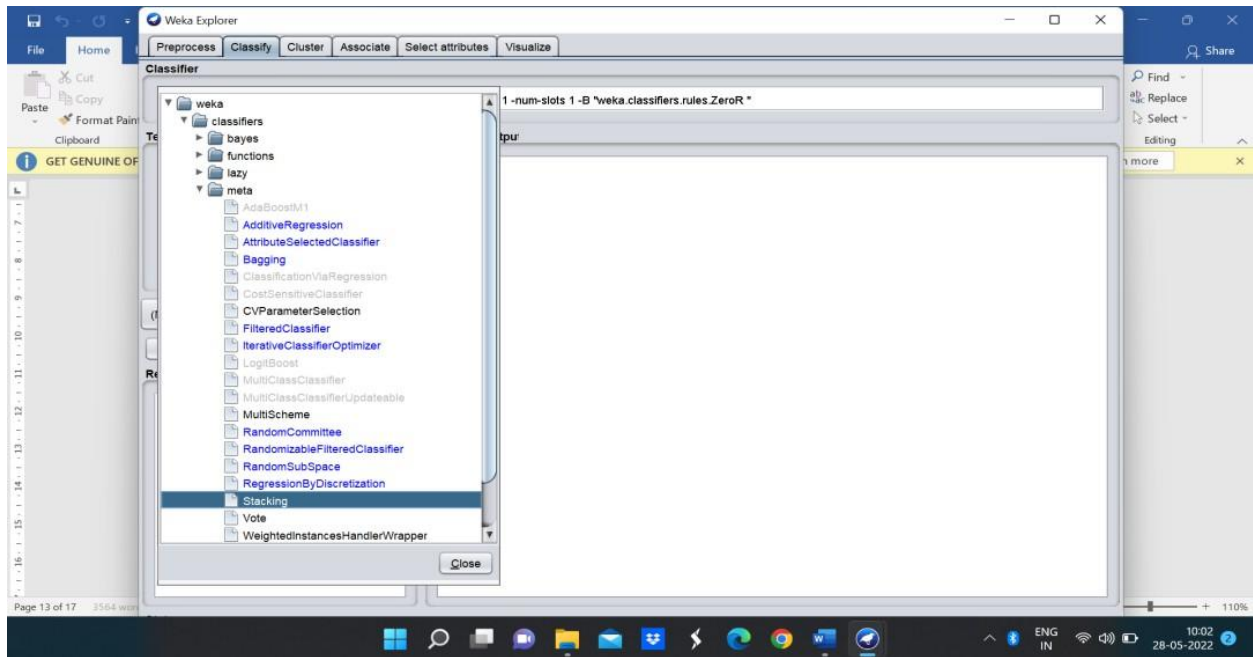


Figure 4: Staking classification using WEKA

The classification techniques employed involve the application of stacking within the WEKA mining tool, with the outcomes presented in table [2].

Table 2: Error presentation using Staking classification

Feature	Value
Correlation coefficient	0
Mean absolute error	8056
Root mean squared error	7.4213
Relative absolute error	100 %
Root relative squared error	100 %
Total Number of Instances	19091

6.2 Meta Data Classification Using Bagging [WEKA Defined]

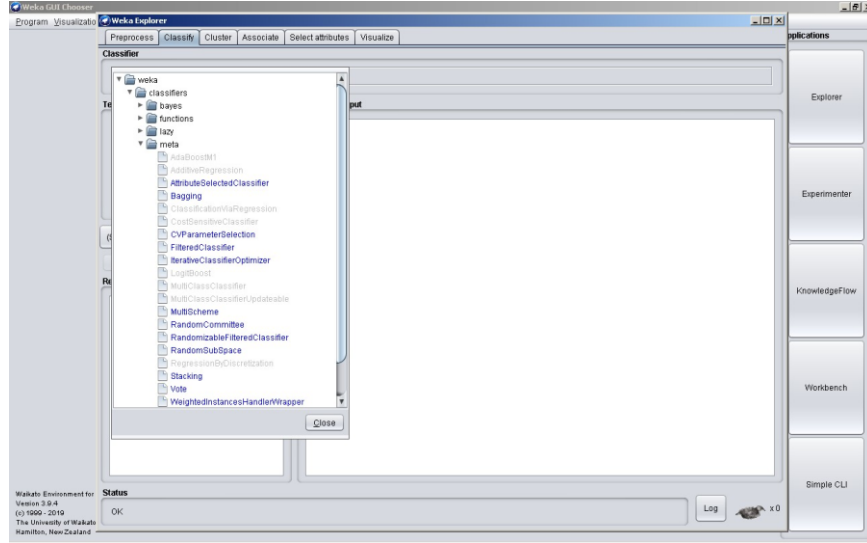


Figure 5: Bagging Classification function under WEKA

Following the implementation of the Staking classification technique for string attributes, as well as the bagging classification approach for numerical data fields, there is a necessity to establish a comparable concept that can be applied to both types of data. Consequently, the execution of a dual stake approach is advantageous, facilitating the mapping of attribute interrelations.

7. A COMPLETE PIPELINE PROPOSED TO DESIGN DATA MODEL PROCESSING

The concept of deep learning is thoroughly articulated, encompassing the data mining process pertinent to our research, as illustrated in figure [6].

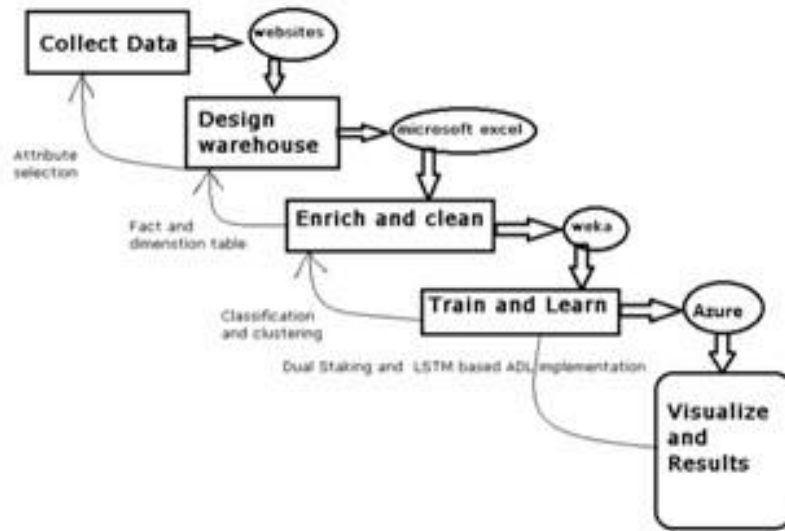


Figure 6: Pipe Line figure of experimental process

8. CONCLUSION

This research presented an innovative methodology for rapid crop yield prediction by combining Dual Stacking with Long Short-Term Memory (LSTM)-based Analytical Deep Learning (ADL). The proposed DSA-LSTM

framework effectively enhances the forecasting accuracy of crop yields, utilizing both short-term and long-term memory variables to accommodate the complexities of agricultural data. By introducing Dual Stacking as a temporary repository for optimized data processing, the model not only improves the speed of yield prediction but also ensures the scalability and adaptability of the system for real-time agricultural decision-making. The integration of set covering optimization techniques allowed for the efficient selection of the most influential attributes, streamlining the data input and refining the overall prediction process. Through the application of data mining tools like WEKA, various classification and clustering techniques were explored, validating the robustness of the proposed methodology. By processing vast amounts of agricultural data from Haryana's Hissar district, the model successfully predicted crop yields with a high degree of accuracy. Additionally, the research highlights the importance of precise data mapping and fact-table management, showing that the accuracy of predictions relies heavily on the quality and organization of the input data. The dual stack mechanism, with its ability to process attributes in parallel, ensures that the system can handle the large-scale and multidimensional data sets common in agricultural applications. Ultimately, this study contributes to the growing body of knowledge in the field of agricultural data science, providing a novel solution for crop yield forecasting that can be readily applied in real-world settings. The proposed methodology not only supports farmers and agricultural stakeholders in making data-driven decisions but also offers a scalable framework that can be adapted to various regions and crop types.

Future Scope

The model can be expanded to incorporate additional environmental, economic, and market factors, further improving its accuracy and utility. Further refinement of the Dual Stacking approach and its integration with other machine learning techniques could also pave the way for even faster and more reliable predictions. As the adoption of precision agriculture grows, this research opens the door to more effective and efficient agricultural practices that are essential for meeting global food security challenges.

REFERENCES

1. Akter, J., Kamruzzaman, M., Hasan, R., Khatoon, R., Farabi, S. F., & Ullah, M. W. (2024, September). Artificial intelligence in American agriculture: a comprehensive review of spatial analysis and precision farming for sustainability. In 2024 IEEE International Conference on Computing, Applications and Systems (COMPAS) (pp. 1-7). IEEE.
2. Atapattu, A. J., Perera, L. K., Nuwarapaksha, T. D., Udumann, S. S., & Dissanayaka, N. S. (2024). Challenges in achieving artificial intelligence in agriculture. In *Artificial intelligence techniques in smart agriculture* (pp. 7-34). Singapore: Springer Nature Singapore.
3. Ganeshkumar, C., Jena, S. K., Sivakumar, A., & Nambirajan, T. (2023). Artificial intelligence in agricultural value chain: review and future directions. *Journal of Agribusiness in Developing and Emerging Economies*, 13(3), 379-398.
4. Hachimi, C. E., Belaiziz, S., Khabba, S., Sebbar, B., Dhiba, D., & Chehbouni, A. (2022). Smart weather data management based on artificial intelligence and big data analytics for precision agriculture. *Agriculture*, 13(1), 95.
5. Linaza, M. T., Posada, J., Bund, J., Eisert, P., Quartulli, M., Döllner, J., ... & Lucat, L. (2021). Data-driven artificial intelligence applications for sustainable precision agriculture. *Agronomy*, 11(6), 1227.
6. Misra, N. N., Dixit, Y., Al-Mallahi, A., Bhullar, M. S., Upadhyay, R., & Martynenko, A. (2020). IoT, big data, and artificial intelligence in agriculture and food industry. *IEEE Internet of things Journal*, 9(9), 6305-6324.
7. Yu, L., Du, Z., Li, X., Zheng, J., Zhao, Q., Wu, H., ... & Huang, X. (2025). Enhancing global agricultural monitoring system for climate-smart agriculture. *Climate Smart Agriculture*, 2(1), 100037.
8. Thiaw, C. M., Asie, L. R., Lovince, H. B., Rousseau, A. N., & Celicourt, P. (2025). A Farm-To-Fork Framework to Assess the Scope and Limitations of Agricultural Data Structures. *Modern Agriculture*, 3(2), e70022.
9. Narang, G., Galdelli, A., Pietrini, R., Solfanelli, F., & Mancini, A. (2024, October). A Data Collection Framework for Precision Agriculture: Addressing Data Gaps and Overlapping Areas with IoT and Artificial Intelligence. In 2024 IEEE International Workshop on Metrology for Agriculture and Forestry (MetroAgriFor) (pp. 580-585). IEEE.
10. Jafar, A., Bibi, N., Naqvi, R. A., Sadeghi-Niaraki, A., & Jeong, D. (2024). Revolutionizing agriculture with artificial intelligence: plant disease detection methods, applications, and their limitations. *Frontiers in Plant Science*, 15, 1356260.
11. Williamson, H. F., Brettschneider, J., Caccamo, M., Davey, R. P., Goble, C., Kersey, P. J., ... & Leonelli, S. (2023). Data management challenges for artificial intelligence in plant and agricultural research. *F1000Research*, 10, 324.
12. Araujo, S. O., Peres, R. S., Ramalho, J. C., Lidon, F., & Barata, J. (2023). Machine learning applications in agriculture: current trends, challenges, and future perspectives. *Agronomy*, 13(12), 2976.
13. Araújo, S. O., Peres, R. S., Barata, J., Lidon, F., & Ramalho, J. C. (2021). Characterising the agriculture 4.0 landscape—emerging trends, challenges and opportunities. *Agronomy*, 11(4), 667.

14. Khanal, S., Kc, K., Fulton, J. P., Shearer, S., & Ozkan, E. (2020). Remote sensing in agriculture—accomplishments, limitations, and opportunities. *Remote sensing*, 12(22), 378
15. Weersink, A., Fraser, E., Pannell, D., Duncan, E., & Rotz, S. (2018). Opportunities and challenges for big data in agricultural and environmental analysis. *Annual Review of Resource Economics*, 10(1), 19-37.
16. Antle, J. M., Basso, B., Conant, R. T., Godfray, H. C. J., Jones, J. W., Herrero, M., ... & Wheeler, T. R. (2017). Towards a new generation of agricultural system data, models and knowledge products: Design and improvement. *Agricultural systems*, 155, 255-268.
17. Janssen, S. J., Porter, C. H., Moore, A. D., Athanasiadis, I. N., Foster, I., Jones, J. W., & Antle, J. M. (2017). Towards a new generation of agricultural system data, models and knowledge products: Information and communication technology. *Agricultural systems*, 155, 200-212.
18. Han, H., Zeeshan, Z., Talpur, B. A., Sadiq, T., Bhatti, U. A., Awwad, E. M., ... & Ghadi, Y. Y. (2024). Studying long term relationship between carbon Emissions, Soil, and climate Change: Insights from a global Earth modeling Framework. *International Journal of Applied Earth Observation and Geoinformation*, 130, 103902.
19. Li, L., Zhang, Y., Wang, B., Feng, P., He, Q., Shi, Y., ... & Yu, Q. (2023). Integrating machine learning and environmental variables to constrain uncertainty in crop yield change projections under climate change. *European Journal of Agronomy*, 149, 126917.
20. Li, L., Wang, B., Feng, P., Li Liu, D., He, Q., Zhang, Y., ... & Yu, Q. (2022). Developing machine learning models with multi-source environmental data to predict wheat yield in China. *Computers and Electronics in Agriculture*, 194, 106790.
21. Williams, K. J., Belbin, L., Austin, M. P., Stein, J. L., & Ferrier, S. (2012). Which environmental variables should I use in my biodiversity model? *International Journal of Geographical Information Science*, 26(11), 2009-2047.