

Article

A Novel Approach in Prediction of Crop Production Using Recurrent Cuckoo Search Optimization Neural Networks

Aghila Rajagopal ¹, Sudan Jha ² , Manju Khari ³, Sultan Ahmad ^{4,*} , Bader Alouffi ⁵  and Abdullah Alharbi ⁶

- ¹ Department of Computer Science and Business Systems, Sethu Institute of Technology, Virudhunagar 626115, India; aghila25481@gmail.com
 - ² School of Sciences, Christ (Deemed to be University), NCR-New Delhi Campus, Ghaziabad 201003, India; jhasudan@ieee.org
 - ³ School of Computer and Systems Sciences, Jawaharlal Nehru University, New Delhi 110067, India; manjukhari@gmail.com
 - ⁴ Department of Computer Science, College of Computer Engineering and Sciences, Prince Sattam Bin Abdulaziz University, Alkharj 11942, Saudi Arabia
 - ⁵ Department of Computer Science, College of Computers and Information Technology, Taif University, P.O. Box 11099, Taif 21944, Saudi Arabia; balouffi@tu.edu.sa
 - ⁶ Department of Information Technology, College of Computers and Information Technology, Taif University, P.O. Box 11099, Taif 21944, Saudi Arabia; amharbi@tu.edu.sa
- * Correspondence: s.alisher@psau.edu.sa

Abstract: Data mining is an information exploration methodology with fascinating and understandable patterns and informative models for vast volumes of data. Agricultural productivity growth is the key to poverty alleviation. However, due to a lack of proper technical guidance in the agriculture field, crop yield differs over different years. Mining techniques were implemented in different applications, such as soil classification, rainfall prediction, and weather forecast, separately. It is proposed that an Artificial Intelligence system can combine the mined extracts of various factors such as soil, rainfall, and crop production to predict the market value to be developed. Smart analysis and a comprehensive prediction model in agriculture helps the farmer to yield the right crops at the right time. The main benefits of the proposed system are as follows: Yielding the right crop at the right time, balancing crop production, economy growth, and planning to reduce crop scarcity. Initially, the database is collected, and the input dataset is preprocessed. Feature selection is carried out followed by feature extraction techniques. The best features were then optimized using the recurrent cuckoo search optimization algorithm, then the optimized output can be given as an input for the process of classification. The classification process is conducted using the Discrete DBN-VGGNet classifier. The performance estimation is made to prove the effectiveness of the proposed scheme.

Keywords: data mining; crop production; recurrent cuckoo search optimization algorithm; Discrete DBN-VGGNet classifier



Citation: Rajagopal, A.; Jha, S.; Khari, M.; Ahmad, S.; Alouffi, B.; Alharbi, A. A Novel Approach in Prediction of Crop Production Using Recurrent Cuckoo Search Optimization Neural Networks. *Appl. Sci.* **2021**, *11*, 9816. <https://doi.org/10.3390/app11219816>

Academic Editor: Amerigo Capria

Received: 18 September 2021

Accepted: 16 October 2021

Published: 20 October 2021

Publisher's Note: MDPI stays neutral with regard to jurisdictional claims in published maps and institutional affiliations.



Copyright: © 2021 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

India's small farms are distinctive. More than 75 percent of the country's overall capital is less than 5 hectares. Most plants are nourished by rain, which only irrigates around 45% of the ground. Approximately 55% of the total Indian population relies on agriculture, according to some figures. In the United States, it is around 5% due to strong agricultural mechanization. India is one of the largest farm products manufacturers, and its productivity remains very limited. Productivity has to be improved so that farmers can gain more with less energy from the same piece of land. Precision farming offers a means of achieving this. Precision agriculture, as it is called, requires the implementation, at the right time, of accurate and accurate overall commentaries such as pee, fertilizers, soil, etc., to improve their productivity and increase their yields [1,2]. The primary income source in India is agriculture. For around two-thirds of India's population, it is their

only source of income [3]. Agricultural sectors accounted for almost 43 percent of India's geographical area. For decades, simple food crops have been produced using farming. Due to the population growth and technological change, changes need to be made to the current cultivation system. The main challenges of existing farming activities are high demand, a full income, and healthy crops. In order to solve this, the focus is therefore on selecting crops that are best adapted to a specific region. Decisions on crop selection include increasing productivity and optimizing the benefits. These cannot be rendered with single parameters; instead, it is necessary to consider further criteria. In order to address several problems in different systems, several decision-making methods were previously developed. None of these have the greatest results, though. Agricultural data mining is a recent area of study. Data mining is the method of identifying, in vast databases, previously discovered and probably remarkable trends. The mined data are part of a model with a systematic data structure; this model can be used with new datasets in order to identify and forecast. Therefore, the farming process itself is complex and farmers change their farm practices naturally based on previous crop output and resource availability. The crop portfolio success depends in part upon water quality, soil-based nutrients (which may rely on previous plot assignments), and crop water requirements. Data mining is a method for collecting and converting secret information from a foundation into a comprehensible framework for further use. This is the computer mechanism by which patterns can be found in vast datasets that contain methods for crossing artificial intelligence, machine learning, statistics, and database systems. The end aim is data mining—and the most commonly used kind of data mining is prediction—and statistical data mining. Many algorithms have been developed over the years to collect information from massive datasets. Among other methods, this problem can be solved with various methods: Grouping, the law of association, clustering, and so on. Here we focus on the methods of classification. Unknown samples are categorized using the information provided by a variety of classified samples. This set is commonly referred to as a training set, because it is generally used to train the execution of the classification technique. The classification task may be considered as a supervised process, in which each instance is part of a class, implied by the importance of a specific target attribute or simply the attributes of a class. Classification routines using a wide range of algorithms can influence the classification of documents, and the specific algorithm used can affect this. A large number of research activities are being carried out for agricultural planning, in which an accurate and precise model for crop yield forecasts, crop classification, soil grading, weather predictions, crop disease prediction, and crop phase classification is being developed. Modeling methods were made for mathematical and machine learning. While crop failures continue to be normal in the implementation of this technique, an advanced computerized crop-forecasting approach solves crop selection issues. Consequently, an effective decision-making methodology is needed for crop collection and cultivation. The key objective of this paper is to apply this decision-making strategy to crop selection. For this aim, the hybrid DBN-VGGNet classifier can be used.

The rest of the paper can be organized as follows. Section 1 depicts the basic introduction to the process of decision making and the crop selection process. The related existing methodologies are depicted in Section 2. The problem definition is depicted in Section 3. The technique used for the crop selection is also defined in Section 3, while the findings and explanation for the study are in Section 4. Section 5 ultimately summarizes the paper.

2. Related Works

Previous research suggested a novel extract approach for the optimized sub-set function [4]. Cultivation depending on the forms to be tracked for algorithm items, dependent on the vector machine, help with classification. The performance provided by these technologies can be observed to have a total accuracy of about 89.6 percentage points.

The authors of suggested an idea focused on fuzzy sets. When fluidity occurs, the brightness amount should be assumed in-pixel in the images for the calculated degree [5].

Then, the ambiguity of the images is treated effectively in the fuzzy package, and IFSs in particular are processed. Following this, the activation of the segmentation can be calculated by the satellite, by which the number of unknown images captured can be decreased. Then, the segmentation of the deficit of the crops for the clustering technique will fuse an image, since it depends on the interval between the intuitionist fuzzy set.

A previous study used a chart rice crop-based neural network algorithm and forecast the yield in the district of Terai [6]. Moreover, the authors of [7] evaluated the production, water output (WEE), output in precipitation usage (PUE), and net crop returns based on proven crop responses to the use of water by integrating the rotation of grass/broadleaf.

Other authors emphasized the degree to which different environmental factors influence rainfall and also used decisions made on crop development, including the identification of diseases and crop selection [8]. Previous authors [9] conducted a comparative study of the GIS-based crop prediction interface machine-learning algorithms (CSPM).

Furthermore, presented a professional knowledge support program for rice, coffee, and cocoa production, focused on the user's input, and external details, such as the position and environment, which will assist selection processes, tracking, surveillance, detection, pest prevention, and selection of fertilizer, among others [10].

In addition, presents the creation of an automated device that analyzes and provides advice to farmers on infected paddy photos [11]. The key aim of designing a system for classifying rice diseases is to simplify the detection and classification of rice diseases [12], which include vector supports and artificial neural networks. The crop forecast takes account of parameters such as precipitation quantity, minimum and maximum temperature, soil type, humidity, and the importance of soil pH. Data were obtained from Maharashtra's agriculture website. Data were broken into nine farming regions.

Previous authors have also defined and organized data according to the relative importance of the key variables that affect the yield of sugarcane and created mathematical models for the prediction of the yield of sugarcane through the use of data mining techniques (DMs) [13]. To that end, the databases of several sucrose mills in Sao Paulo, Brazil, were analyzed using three DM techniques. Meteorological variables and plant management have been studied using the following approaches of DM: Forest random, raise, and vector assistance. The resulting models have been evaluated using an independent dataset. As a comparison, a random forest algorithm was used.

The application of machine learning techniques in agriculture to the study of soil fertility [14] was discussed. Agriculture has long been one of the research fields of concern. This research is aimed at evaluating, classifying, and enhancing soil data according to different factors. Technical developments such as robotics and data processing have taken advantage of agricultural science. Data mining nowadays is used in large fields, and many off-shelf database mining products and domain-specific data mining applications sell software, but data mining is a comparatively new field of study in agricultural soil datasets. The vast volumes of data now virtually obtained in connection with plants can be processed and used. The authors of [15] stated that, by evaluating all these challenges and problems, such as atmosphere, temperature, moisture, snow, and moisture, there is no correct way of fixing the situation faced by us in terms of technology. Indian agriculture has various ways to improve economic development.

k-Nearest Nearest-Neighbor Algorithm Potentials, an ET0 estimation data mining tool, have been explored in semi-arid China, using minimal climate data [16]. Furthermore, the PM-56 equation was checked with a KNN-dependent ET0 prediction model.

The new activation function and revised random weight and bias values for crop yield estimation [17], obtained by using various weather-parameter datasets, are used to build an improved MLP neural network. Established activation features and newly developed activation functions, including weights and bias value, have evaluated the MLP model. This research analyzes the outcomes of various activation capabilities and proposes several new basic activation functions such as DharaSig, DharaSigm, and SHBSig, in order to increase the efficiency and accuracy of neural networks. In addition, the DharaSig

functions DharaSig1, DHaraSig2, and DharaSig3 were created by three new activation functions with minor variations.

Another work developed relevant operating laws and the required weighted aggregation operators [18]. In this, the characteristics of the neutral type in the membership degrees of the group and the sum of probability are specified by a neutral addition and scalar multiplication operational rules. Any aspects of the legislations introduced are analyzed.

In [19], the authors discussed the issue of a robust exponential passivity analysis for uncertain neutral-type neural networks with mixed-interval time-varying delays. In another work [20], the Levenberg Marquardt method was used to analyze and predict human gait, and the same method can be used in forest and farming surveys. Traditional surveys are expensive, time-consuming, laborious, and difficult to perform, especially in mountainous areas and dense forests and fields.

Failures in crops are fairly common. At the same time, agriculture is impacted by other causes, such as agricultural degradation, overused fertilizers and insecticides, stresses on toxic substances and radiation, etc. Many bugs have been shown to be resistant to insecticides. Early plant prediction will solve crop production issues. Hence, an appropriate technique for decision making is required for the collection and cultivation process of crops.

3. Proposed Work

The proposed technique based on decision making and the optimization-based selection of crops to improve the crop yield is depicted in Figure 1.

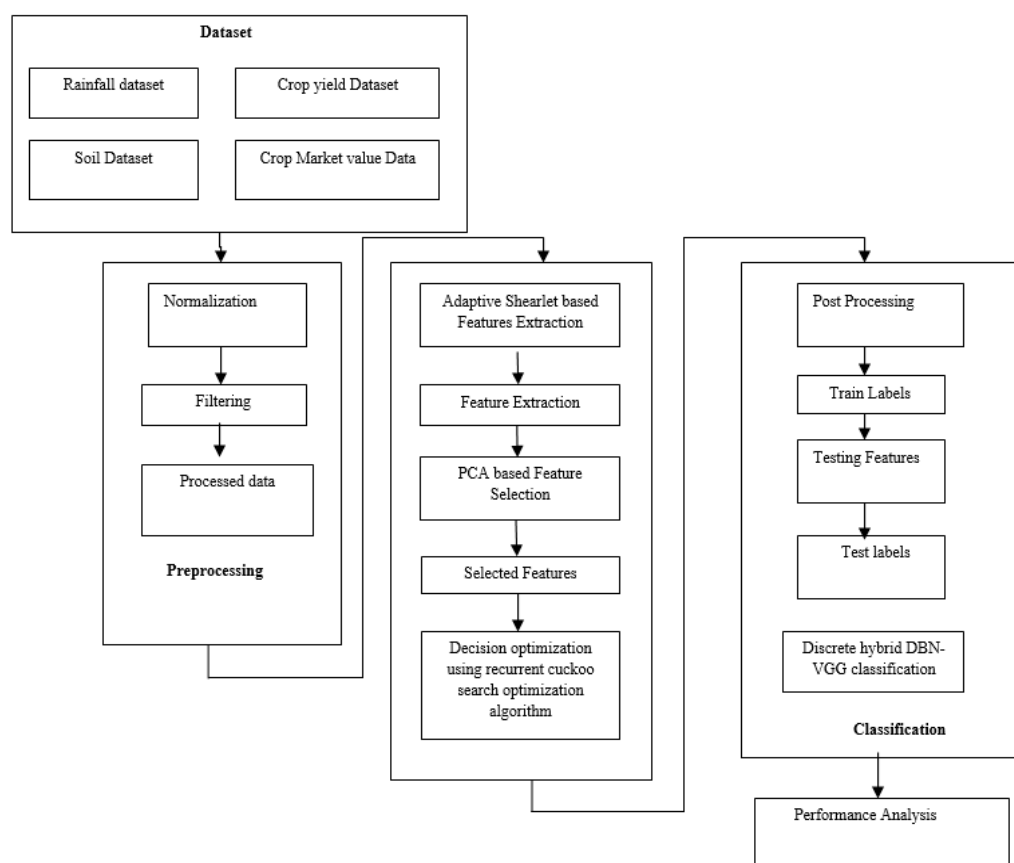


Figure 1. Schematic representation of the proposed methodology.

3.1. Preprocessing

For the process of data processing, there is a need to normalize the values. Some types of the normalization only involve a process of rescaling in order to obtain values related to some other variable. When crop population parameters are known, we need to modify

the errors by performing certain simple adjustments. After modifying the errors, the population values can be normally distributed rather than of a random distribution. The first step of the process of normalization is to obtain the z-score. The z-score is represented in Equation (1).

$$z = [(x - \mu)/\sigma] \quad (1)$$

where μ is the mean of the crop population and σ is the standard deviation of the crop population. When the population mean and the crop population standard deviation are not known, then the standard score will be calculated using the sample mean and sample standard deviation, which can be represented in Equation (2).

$$z = \frac{x - \bar{x}}{s_x} \quad (2)$$

where $\bar{x}$ is the mean of the sample and s_x is the standard deviation of the sample.

The initial stage of data analysis is the pre-processing phase. The data will be translated to a numerical format that is simple to grasp. A significant goal of this method is to remove invaluable terms. The standardization would streamline the pre-processing level. The first phase in the method to assign the integer value is mathematical encoding.

$$J = [(k - \mu)/\sigma] \quad (3)$$

where μ is the mean of the data amount and σ is the standard deviation of the data. When the data mean and standard deviation are not known, the standard integer value will be assigned using the sample mean and standard deviation.

$$J = \frac{k - \bar{k}}{s_x} \quad (4)$$

$\bar{k}$ is the mean of the sample and s_x is the standard deviation of the sample. For proper assignment of the variable, the sequence can be converted into the matrix format.

$$D = k * (k^T k)^{-1} k^T \quad (5)$$

The variance for the matrix is,

$$Var(\hat{v}_i) = \sigma^2(1 - v_{ii}) \quad (6)$$

$$Var(\hat{v}_i) = \sigma^2(1 - \frac{1}{i} - [(k_i - k^2) / \sum_{j=1}^i (k_j - \bar{k}^2)]) \quad (7)$$

Then, the residual which can be calculated by

$$\frac{\hat{v}_i}{\bar{\sigma}} \sqrt{1 - v_{ii}} \quad (8)$$

where $\bar{\sigma}$ is an estimate of σ .

The function scaling method can then be implemented to allocate values for all variables from 0. This is referred to as ordinal or integer encoding

$$Z' = \frac{(J - d_{min} Var(\hat{d}_i))k_i}{(d_{max} - d_{min})\bar{\sigma}^2} \quad (9)$$

3.2. Feature Extraction

Then, the features were chosen using the adaptive Shearlet approach (ASA) following the segmentation step. This is a means of removing second-order mathematical texture characteristics. This technique has been utilized in several applications, whereas the interaction of three or more data attributes occurs with higher-order features. This is a

mathematical task that will typically efficiently eliminate unwanted data. The accuracy of the data may also be rendered clearly. The data becomes differentiated during the analysis cycle. In a particular exact differential feature, Shearlet can specify the frequency of the data attributes. The single pieces of data are questioned here, and another pixel is known as the $\emptyset$ route l and the adjacent value detachment of m . Usually m obtains a single value, and $\emptyset$ can benefit directionally. The obtained directional value can remove the attributes of the data used for the classification process. The Shearlet process may be set as follows:

$$K(l, n) = G(l, n, o, \emptyset) / \sum_{l=1}^H \sum_{n=1}^H G(l, n, o, \emptyset) \quad (10)$$

where G is the frequency vector l, n, o , the frequency of the particular component will generally having the pixel values of 0 and 1, K represents the features of data, (l, n) is the component of l and n , and $\emptyset$ represents the normalized constant.

ASA is performed with N -dimensional vectors on the database representing the position of individual components inside the feature space. The characteristics may then be chosen according to their correlation. The approach assesses the subset by taking into consideration each function's predictive potential and redundancies (or similarity) individually. This implies that the algorithm will decide its next steps, provided by a (heuristic) function, by choosing the option that optimizes the performance of this feature. Heuristic functions may also be built to reduce the expense of the target. By using the following equation, the correlation between the features can be defined.

$$K\left(\frac{o}{\partial}, \mu\right) = \left[\frac{\varphi(\partial + \mu)}{\varphi(\partial)\varphi(\mu)} \right] \delta(\partial + \mu) \gamma(\mu - 1) \quad (11)$$

Afterwards, some of the important crop features that can be extracted are depicted below.

$$\text{Minimal capital} = \frac{1}{l} - 1 \sum_{l=1}^{l-1} a(j+1) - y_i(j) \quad (12)$$

$$\text{High yield} = \sum_{i,j=0}^{n-1} F(i, j) \left[\frac{(i - \mu i)(j - \mu j)}{\sqrt{(\sigma i^2)} \sqrt{(\sigma j)^2}} \right] \quad (13)$$

$$\text{Flexible marketing} = \sum_{i,j=0}^{n-1} \frac{F(i, j)}{F} - (F + 2) \quad (14)$$

This method helps to process the data and extract the features of a crop from the data in an effective manner.

3.3. Feature Selection

In order to select the appropriate farm crops to be cultivated in the defined experimental field, a decision model has been created. The 26 selected and grouped input variables are sorted into 6 primary variables, namely land, water, season, input, profit, and yield. In the course of extracting features, 26 input variables were gathered in order to build the decision-making model to help farmers make decisions on crop selection at the respective agricultural locations. Six of the twenty-six input variables are selected using PCA (Principal Component Analysis) by utilizing the feature selection method. Then the features can be visualized using the PCA system. An orthogonal transformation-driven function design methodology is the primary component of the system. The number of key elements is fewer than the original criteria. Input values can be limited by PCA. The parameter value of every class can be shown in this PCA function map. PCA is a technique for reducing dataset dimensionality, by increasing interpretability. PCA condenses information from a wide variety of variables into fewer variables by introducing some kind of transformation theory. Correlation implies the information is repetitive and that if this is consistent, information

may be compact. Considering the two $F1$ and $F2$ features, these are distributed uniformly on a $[-1, 1]$ binary and the O output class, which are given below.

$$O = \begin{cases} 0 & \text{if } F1 + F2 < 0 \\ 1 & \text{if } F1 + F2 \geq 0 \end{cases} \quad (15)$$

In this dilemma, the data points are provided in the shady regions. The problem is linearly separable, and the required features of $F1 + F2$ can easily be chosen. The LDA is conducted with N -dimensional vectors on the collection of data providers, indicating the path of the function space. This vector provides the best information about the problem and provides the latest function to the output class that projects it into the space $(F1, F2)$. For each and every class, the matrix will be derived between the S_{kr} class and within the S_{ps} class scatter matrix, which are defined as follows,

$$S_{ps} = \sum_{a=1}^{ps} S_a; S_a = \frac{1}{p_s} \sum_{p \in p_s} (p - l_a)(p - l_a)^T \quad (16)$$

$$S_{kr} = \sum_{a=1}^{kr} (l_a - l)(l_a - l)^T \quad (17)$$

where $d \times d$ is matrix A , which is used for dimensionality reduction to ensure d dimensional features $y = ATx$. The covariance matrix of all samples is given by,

$$P = \frac{1}{n} \sum_{p \in P} (p - l)(p - l)^T \quad (18)$$

By obtaining the covariance matrix, the pointed features can be sorted.

Finally, by implementing PCA, important variables such as land, water, season, input, profit, and yield are selected.

The obtained results are again optimized to improve the process further. Here the optimization was conducted using the improved chicken swarm optimization algorithm. The frameworks of the chicken swarm are first reviewed before expanding the chicken swarm optimization into multi-target problems. The new bioinspired algorithm for the single goal optimization is Recurrent Chicken Swarm Optimization (RCSO). RCSO simulates the hierarchical order and actions of a chicken swarm in the food quest, where any chicken receives an optimization problem as a possible solution. Essentially, CSO uses the four laws to idealize chickens' actions:

1. The chicken swarm includes multiple groups, including the main rooster, several hens, and several chicks together. There are several categories.
2. Within the chicken swarm, each category relies on the fitness value of the individual, namely the individual identification (roosters, hens, and chicks). The best fitness values for the chickens are known as roosters: Each one is a rooster in a group. Chickens with lower health are known as chicks. The rest are hens.
3. The hierarchical hierarchy, supremacy, and the mother-child relationship will shift entirely after increasing the number (G) of moves.
4. The chickens are searching for food after their rooster friend.

We suspect the chickens would inadvertently eat the great food that many have already discovered. Chicks appeal to their mothers for food. In the food rivalry, the rooster has the edge. We see RN, HN, CN, and MN in a chicken swarm of n people, reflecting the number of roosters, hens, chicks, and mother hens, respectively. Each chicken is shown in a sizeable space by its position. CSO makes three kinds of chicken and each type has its own calculation movement. Chickens will move in 1-D, 2-D, 3-D, and super dimensional space with modified position vectors. In a population-based ICSO algorithm, the n chicken swarm is used as the search agent in the field of question. The implementation of this algorithm determines the optimal solution for crop cultivation. Then, the chicken searches for a suitable way to upgrade it. Roosters with the highest fitness values will find food in

a wider number of locations than roosters with worse fitness values. They are defined in Equation (24),

$$N(R_i, P_j) = N_i \rightarrow e^{bt} \cdot F \cos(2\pi t) + P_j, \quad (19)$$

where N_i is a Gaussian distribution, P is a rooster index, which is randomly selected from the roosters' group, and ϕ is the fitness value of the rooster P .

Hens' actions oppose their nurturing partner. In comparison, the good food that other chickens discover, while repressed by the other chickens, hens will occasionally rob. A more aggressive hen will contend for food with the more submissive. Mathematically, the movement of hens can be formulated as in

$$\sigma_{km} = \frac{G_{km}}{G_{km(\max 1)}} = \frac{\sqrt{(G_{kz} - G_{mz})^2 + (G_{ky} - G_{my})^2}}{|G_k| * |G_m|} \quad (20)$$

where σ_{km} is the movement of the hens. The hierarchic relationship upgrade provides chicks the ability to become roosters for the right solutions (i.e., it gives further urgency to find the best solutions for chicks). The suitable crop in each iteration after adjusting the crop list i classified according to health values. The number of crops to pursue is minimized in relation to the iteration and, to a large degree, enables the most promising alternatives to be used—the ICSO is used here primarily to determine the best potential health for the further extraction method. The optimal benefit will be determined upon completion of the optimization.

$$M_{fitness} = (M - 1 * (M - 1/T)) \quad (21)$$

where l is the iteration number, M is the number of optimal paths, and T is the maximum number of iterations.

After that, the best fitness value can be calculated using optimization. The maximum fitness concentration value can be found along with its corresponding position. For each optimization algorithm, the fitness value can be calculated.

$$Best = \min_{i=1}^M OS_*^i \quad (22)$$

where OS is the optimal solution for crop cultivation having an i th iteration number.

3.4. Discrete Hybrid DBN-VGG Classification

The DBN network can be refined as neural. DBN has several non-linear hidden layers, it can be pretrained to serve as a non-linear reduction of the dimensionality of the inputs, and the network trainer can be further sensory feedback. Similarly, VGG is a K-proposed variant of the neural network. ImageNet, which consists of over 14 million images from 1000 different grades, achieved 92.7 percent top-5 test precision. This was one of the most famous models of ILSVRC-2014. It improves Alex Net through the replacement of broad-kernel filters (11 and 5 in the first and second coevolutionary strata) with multiple three-kernel-sized filters one by one. VGG16 was trained for weeks and using NVIDIA Titan Black GPU. Here, the deep belief neural network can be merged with VGG to form a hybrid for an effective abnormality classifying purpose. Therefore, the disease must be distinguished as to whether it is suitable for the particular area, depending upon the characteristics to be extracted and optimized. To enhance the extraction and recovery process, the classification method is used to distinguish retinal fundus images according to their different characteristics. The Discrete hybrid DBN-VGG classification is used to derive the non-linear properties of the retinal atherosclerosis fundus images for better or more specific functionality. The application of both DBN-VGG classification methods demonstrates the fundus images with their current precision and efficiency and enhanced consistency. The featured images were used to identify particular crops. The method of classification relies on the characteristics derived. The DBN-VGG measures the likelihood and executes a task. There is an overall distribution. The data are first interpreted and

resized by classifiers throughout the method, and then the method of classification is conducted by measuring the class probability.

$$\text{classify}(F) = a_j^l b_j^l \quad (23)$$

The classification is concluded as

$$F = a_j^l b_j^l - n \left(a_j^l b_j^l \right)^2 \quad (24)$$

where F is the feature, n is the pointed feature, and $a_j^l b_j^l$ are the classified features.

The Pseudocode of classification process can be shown as below in Algorithm 1.

Algorithm 1 Pseudocode for the Classification Process

Input: Enhanced data
Output: filtered data
Initialize the DBN layers
Initialize the vgg layers
Integrated DBN layers and vgg layers
Initialize train features
Initialize label
Train label = 80%
Test label = 20%
Lab = unique(label)
For ii = 1:length(Lab)
Class = find(label == Lab(ii))
To pass trainfeatures data to each layers,
Traincut = length(class) – traincut
Traindata = [traindata; trainfeatures; class(1: Traincut)end-5:end]
Predict label = classify(net,traindata)
End
End
For ii = 1:size(traindata,1)
Traindata = [traindata; trainfeatures; class(1: Traincut)end-5:end]
End
For ii = 1:size(trainfea,1)
Testdata = [trainfea; trainfeatures; class(1: Traincut)end-5:end]
End

4. Performance Analysis

This section examines the experimental results of the implemented model. The new decision-making model can be implemented here, and simulation can be carried out. The tests would have higher accuracy, and the consistency would be more dependable. Therefore, the overall success of the approach proposed in this section shall be measured. The outcomes and viability of the suggested solution are calculated and contrasted with the parameters of the calculation.

4.1. Performance Metrics

Accuracy

This is a plot of organized errors, a predisposition measure; poor accuracy produces a difference between an outcome and a true value. In certain instances, data are checked with the same method, and the exact performance of the implemented model is evaluated. The quality of the overall data is the percentage of actual outcomes (both positive and negative).

$$\text{Accuracy (A)} = (TP + TN) / (TP + TN + FP + FN) \quad (25)$$

Precision

Precision is a portrayal of random errors that is a measure of algebraic variability.

$$\text{Precision} = \text{TP} / (\text{TP} + \text{FP}) \quad (26)$$

Recall

Recall in certain fields calculates the proportion of positive facts and correctly defines the true optimistic rate, warning, or probability of identification.

$$\text{Sensitivity} = \text{TP} / (\text{TP} + \text{FN}) \quad (27)$$

Mean square error

This measures the average of the squares of the errors, that is, the average squared difference between the estimated values and the actual value.

$$\text{MSE} = \frac{1}{N} \sum_{i=1}^N (x_i - \hat{x}_i)^2 \quad (28)$$

Dice co-efficient:

We assume b and c are the properties of the fundamental truth data and the data characteristics found. Then, we can then determine the Dice coefficient as

$$D(b,c) = (2b \cap c / a + b) = 2\text{TP} / 2\text{TP} + \text{FN} + \text{FP} \quad (29)$$

Jaccard co-efficient:

The resemblance of the two groups can be calculated as

$$J(b,c) = (b \cap c / b \cup c) = (b \cap c) / b + c - (b \cap c) = (\text{TP} / \text{TP} + \text{FN} + \text{FP}) \quad (30)$$

Regarding the solution, Figure 2a,b is a dice and Jaccard ranking. The findings reveal that the proposed method shows exact, high Jaccard values and the coefficient of dices.

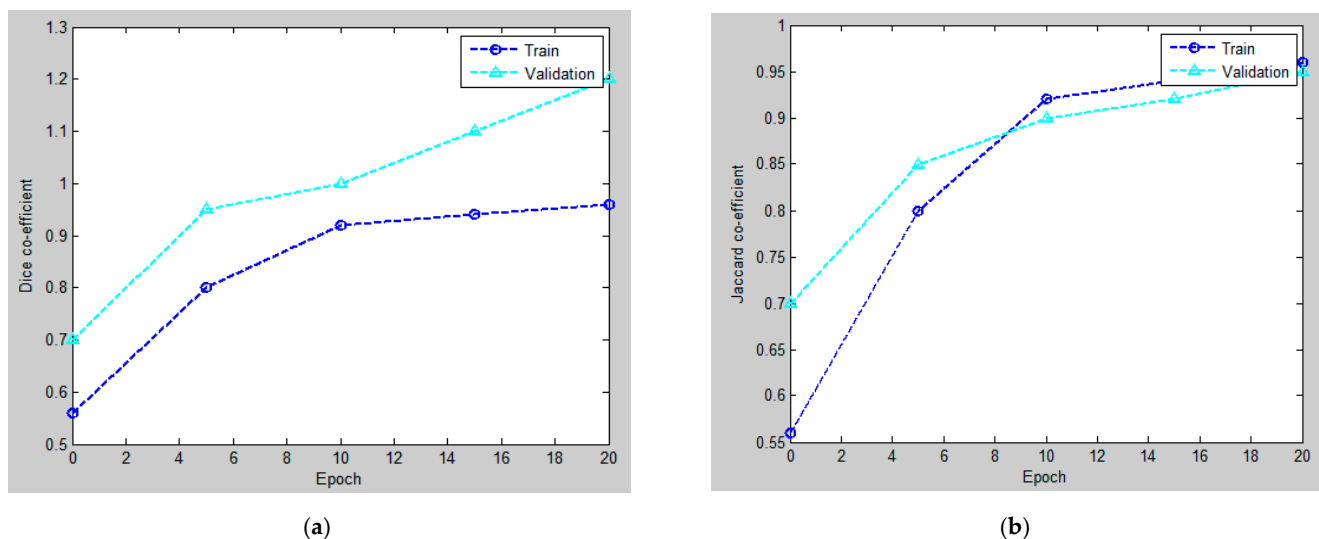


Figure 2. Proposed dice (a) and Jaccard (b) comparative analysis.

4.2. Comparative Performance Analysis

In this section, certain novel methodologies used in existing works [21,22] are briefly introduced.

The idea is contrasted with other algorithms and methods for identification, though based on significant performance metrics. Due to the increased success compared to other, existing methods, the Discrete Hybrid DBN-VGG + RCSO was used as the crop selection

method. Figure 3 represents the dataset of graphical values of the proposed decision-making model output. With the use of the suggested classifier model, the best performance possible is achieved. The findings suggest an improved performance of the proposed method. This performance reflected 97% precision and 97% accuracy with 94% recall. The overall efficiency of the proposed model is depicted in Figure 3 and Table 1.

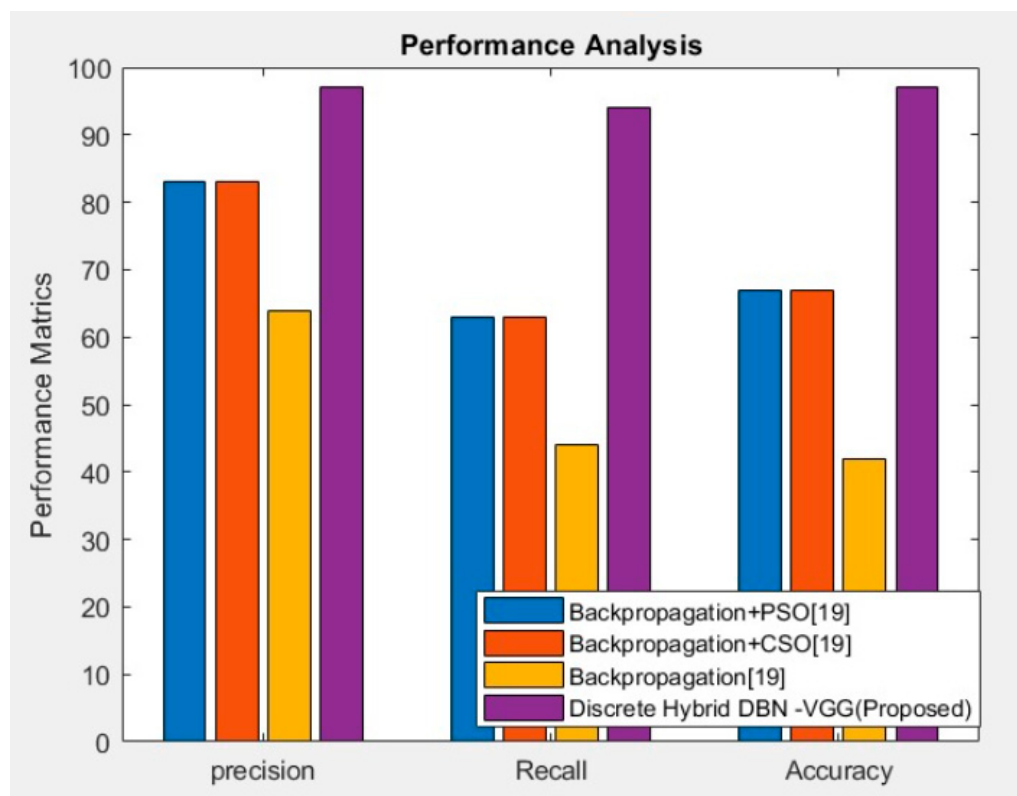


Figure 3. Comparative analysis of proposed and existing techniques.

Table 1. Comparative performance analysis.

Parameters	Backpropagation + PSO [21]	Backpropagation + CSO [21]	Backpropagation [21]	Discrete Hybrid DBN-VGG RCO (Proposed)
Precision	83%	83%	64%	97%
Recall	63%	63%	44%	94%
Accuracy	67%	67%	42%	97%
MSE	0%	0.10%	0.15%	0.01%

The AROC curve represents the sensitivity and specificity performance. Figure 4 represents the performance of AROC, which is 0.9677. Here, from Figure 4, the value of AROC is 0.9677, which indicates the classifier distinguishes the class perfectly. From the results obtained, it is clear that the proposed method outperforms other existing methods.

By using the decision optimization model and Discrete Hybrid DBN-VGG Classification, the best crop was selected, depending upon the rank obtained by the crop as shown in Table 1. When compared to others, wheat, rice, and millets have a higher probability of achieving high yield than others.

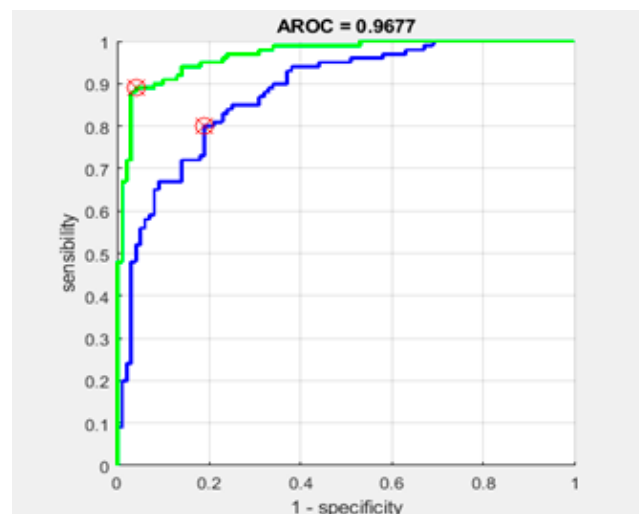


Figure 4. Performance of AROC.

5. Conclusions

In this paper, we use a Discrete Hybrid DBN-VGG + RCSO algorithm that can help to classify the variety of crop-based planting schedules. From this application, the proposed model focuses on a theoretical model that can be attained for the performance of better separation among the various types of crops in a planting-based schedule. Here, for assessing the effective performance of the implemented method, it was compared with three other existing methods that have been proposed very recently. The performance of the proposed method shows effective results when compared to other existing methods. This clearly shows that the proposed approach can select the crop that can yield higher profit in an effective manner.

Author Contributions: Conceptualization, A.R., M.K., S.J. and S.A.; methodology, M.K., A.R. and A.A.; software, A.R. and M.K.; validation, A.R., B.A., S.J. and S.A.; formal analysis, A.R. and A.A.; investigation, A.R., B.A. and A.A.; resources, A.R., S.J. and S.A.; data curation, A.R., B.A. and A.A.; writing—original draft preparation, A.R., S.J. and S.A.; writing—review and editing, B.A., A.A. and S.A.; visualization, A.R., S.A. and S.J.; supervision, M.K.; project administration, A.R., A.A. and M.K.; funding acquisition, A.R. and A.A. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Acknowledgments: This research was supported by the Taif University Researchers Supporting Project number (TURSP-2020/314), Taif University, Taif, Saudi Arabia.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Messina, C.D.; Technow, F.; Tang, T.; Totir, R.; Gho, C.; Cooper, M. Leveraging biological insight and environmental variation to improve phenotypic prediction: Integrating crop growth models (CGM) with whole genome prediction (WGP). *Eur. J. Agron.* **2018**, *100*, 151–162. [\[CrossRef\]](#)
2. Nosratabadi, S.; Karoly, S.; Beszedes, B.; Felde, I.; Ardabili, S.; Mosavi, A. Comparative Analysis of ANN-ICA and ANN-GWO for Crop Yield Prediction. In Proceedings of the 2020 RIVF International Conference on Computing and Communication Technologies (RIVF), Ho Chi Minh City, Vietnam, 14–15 October 2020.
3. Kumar, D.M.; Kumar, K.M. Decision Support System for Prediction of Crop Yield Using Re-Engineered Artificial Neural Network. *J. Crit. Rev.* **2020**, *7*, 2290–2297.
4. Cui, J.; Zhang, X.; Wang, W.; Wang, L. Integration of optical and SAR remote sensing images for crop-type mapping based on a novel object-oriented feature selection method. *Int. J. Agric. Biol. Eng.* **2020**, *13*, 178–190. [\[CrossRef\]](#)
5. Ananthi, V. Fused Segmentation Algorithm for the Detection of Nutrient Deficiency in Crops Using SAR Images. In *Artificial Intelligence Techniques for Satellite Image Analysis*; Springer: Berlin/Heidelberg, Germany, 2020; pp. 137–159.

6. Baidar, T. Rice Crop Classification and Yield Estimation Using Multi-Temporal Sentinel-2 Data: A Case Study of Terai Districts of Nepal. 2020. Available online: <https://www.semanticscholar.org/paper/Rice-crop-classification-and-yield-estimation-using-Baidar/bf0c1e646caf560db7d19567c8ee10fec0ec3372> (accessed on 21 February 2020).
7. Nielsen, D.C.; Vigil, M.F.; Benjamin, J.G. Evaluating decision rules for dryland rotation crop selection. *Field Crops Res.* **2011**, *120*, 254–261. [\[CrossRef\]](#)
8. Lavanya, K.; Jain, A.V.; Jain, H.V. Optimization and Decision-Making in Relation to Rainfall for Crop Management Techniques. In *Information Systems Design and Intelligent Applications*; Springer: Berlin/Heidelberg, Germany, 2019; pp. 255–266.
9. Tamsekar, P.; Deshmukh, N.; Bhalechandra, P.; Kulkarni, G.; Hambarde, K.; Husen, S. Comparative Analysis of Supervised Machine Learning Algorithms for GIS-Based Crop Selection Prediction Model. In *Computing and Network Sustainability*; Springer: Berlin/Heidelberg, Germany, 2019; pp. 309–314.
10. Lagos-Ortiz, K.; Medina-Moreira, J.; Alarcón-Salvatierra, A.; Morán, M.F.; del Cioppo-Morstadt, J.; Valencia-García, R. Decision Support System for the Control and Monitoring of Crops. In Proceedings of the 2nd International Conference on ICTs in Agronomy and Environment, Guayaquil, Ecuador, 22–25 January 2019; pp. 20–28.
11. Das, S.; Sengupta, S. Feature Extraction and Disease Prediction from Paddy Crops Using Data Mining Techniques. In *Computational Intelligence in Pattern Recognition*; Springer: Berlin/Heidelberg, Germany, 2020; pp. 155–163.
12. Fegade, T.K.; Pawar, B. Crop Prediction Using Artificial Neural Network and Support Vector Machine. In *Data Management, Analytics and Innovation*; Springer: Berlin/Heidelberg, Germany, 2020; pp. 311–324.
13. Hammer, R.G.; Sentelhas, P.C.; Mariano, J.C.Q. Sugarcane Yield Prediction through Data Mining and Crop Simulation Models. *Sugar Tech.* **2020**, *22*, 216–225. [\[CrossRef\]](#)
14. Chaudhari, R.; Chaudhari, S.; Shaikh, A.; Chiloba, R.; Khadtare, T. Soil Fertility Prediction Using Data Mining Techniques. *Mukt. Shabd. J.* **2020**, *9*.
15. Champaneri, M.; Chachpara, D.; Chandvidkar, C.; Rathod, M. Crop Yield Prediction Using Machine Learning. Available online: https://www.researchgate.net/publication/340594772_CROP_YIELD_PREDICTION_USING_MACHINE_LEARNING (accessed on 4 April 2020).
16. Feng, K.; Tian, J. Forecasting reference evapotranspiration using data mining and limited climatic data. *Eur. J. Remote Sens.* **2021**, *54*, 363–371. [\[CrossRef\]](#)
17. Bhojani, S.H.; Bhatt, N.J.N.C. Wheat crop yield prediction using new activation functions in neural network. *Neural Comput. Appl.* **2020**, *32*, 13941–13951. [\[CrossRef\]](#)
18. Garg, H. Neutrality operations-based Pythagorean fuzzy aggregation operators and its applications to multiple attribute group decision-making process. *J. Ambient Intell. Humaniz. Comput.* **2020**, *11*, 3021–3041. [\[CrossRef\]](#)
19. Samorn, N.; Yotha, N.; Srisilp, P.; Mukdasai, K. LMI-Based Results on Robust Exponential Passivity of Uncertain Neutral-Type Neural Networks with Mixed Interval Time-Varying Delays via the Reciprocally Convex Combination Technique. *Computation* **2021**, *9*, 70. [\[CrossRef\]](#)
20. Alharbi, A.; Equbal, K.; Ahmad, S.; Rahman, H.U.; Alyami, H. Human Gait Analysis and Prediction Using the Levenberg-Marquardt Method. *J. Healthc. Eng.* **2021**, *18*, 2021.
21. Layona, R.; Suharjito, S.; Tunardi, Y.; Tanoto, D. Optimization of land suitability for food crops using neural network and swarm optimization algorithm. *Int. Rev. Comput. Softw.* **2016**, *11*, 1. [\[CrossRef\]](#)
22. Zhong, L.; Hu, L.; Zhou, H. Deep learning based multi-temporal crop classification. *Remote Sens. Environ.* **2019**, *221*, 430–443. [\[CrossRef\]](#)