

Classification of Medicinal Plant Leaves using Deep Learning Algorithms

Aruna S. K.*

Department of AI and Data Science Engineering, CHRIST University, Bangalore, India

E-mail: sksaruna@yahoo.co.in

ORCID iD: <https://orcid.org/0000-0002-6638-2772>

*Corresponding author

Praveen P.

General Manager, Sun Publications, Sivakasi, India

E-mail: praveenabhinav7890@gmail.com

ORCID iD: <https://orcid.org/0000-0002-4987-7896>

Gowtham K.

Product Solution Engineer, Juspay Technologies Private Limited, Bangalore, India

E-mail: gowthamks266@gmail.com

ORCID iD: <https://orcid.org/0009-0004-9000-7122>

Mohammed Khashif S.

Software Engineer, Impelsys Private Limited, Bangalore, India

E-mail: mohammedkhashif02@gmail.com

ORCID iD: <https://orcid.org/0009-0007-2126-8355>

Keerthana Jaganathan

School of Computer Science, Northumbria University, Newcastle upon Tyne, NE1 8ST, England, UK

E-mail: keerthana.jaganathan@northumbria.ac.uk

ORCID iD: <https://orcid.org/0009-0004-1413-9377>

K. Karthick

Department of Electrical and Electronics Engineering, GMR Institute of Technology, Rajam, Andhra Pradesh, India

E-mail: karthick.k@gmr.it.edu.in

ORCID iD: <https://orcid.org/0000-0001-7755-5715>

Received: 15 July 2025; Revised: 21 September 2025; Accepted: 26 November 2025; Published: 08 February 2026

Abstract: This research explores the automated leaf-based identification of medicinal plants, utilizing machine learning and deep learning techniques to address the crucial need for efficient plant classification. Driven by the vast potential of medicinal plants in pharmaceutical development and healthcare, the study aims to surpass the limitations of existing methodologies through thorough experimentation and comparative analysis. The primary goal is to develop a robust and automated solution for classifying medicinal plants based on leaf morphology. The methodology encompasses acquiring diverse datasets. Specifically, set 1 data is processed by applying resizing, rescaling, saturation adjustment, and noise removal, while Set 2 data is processed by applying resizing, rescaling, saturation adjustment, noise removal, and PCA (Principal Component Analysis). The proposed algorithms include Support Vector Machines (SVM), Convolutional Neural Networks (CNNs), YOLOv8, Vision Transformer (ViT), ResNet, and Artificial Neural Networks (ANN). The study evaluates the efficacy and effectiveness of each algorithm in plant classification using metrics such as accuracy, recall, precision, and F1 score. Notably, the ResNet model achieved 93.8% and 94.8% accuracy in Set 1 and Set 2, respectively. The SVM model demonstrated 56.5% and 56.6% accuracy in Set 1 and Set 2, while the Vision Transformer (ViT) model achieved 84.9% and 74.4% accuracy in Set 1 and Set 2, respectively. The CNN model showcased high accuracy at 96.7% and 94.8% in Set 1 and Set 2, followed closely by the ANN model with 96.7% and 96.6% accuracy. Lastly, the YOLOv8 model achieved 96.0% and 95.1% accuracy in Set 1 and Set 2, respectively. The comparative analysis identifies CNN and ANN as the top-performing algorithms. This research

significantly contributes to the advancement of medicinal plant identification, pharmaceutical research, and environmental conservation efforts, emphasizing the potential of deep learning techniques in addressing complex classification tasks.

Index Terms: Machine Learning, Deep Learning, Vision Transformer, CNN, Classification.

1. Introduction

In the realm of botanical science, the identification of medicinal plants stands as a fundamental endeavor with far-reaching implications for human health and well-being. With approximately 44,500 identified and classified plant species in India alone, the need for systematic and efficient methods of plant identification becomes increasingly imperative. Leveraging advancements in artificial intelligence and machine learning, this research endeavors to contribute to the automation of plant identification processes, particularly focusing on the classification of medicinal plants based on leaf morphology. Given the vast diversity of plant species and the intricate nuances of their morphological features, traditional methods of plant identification often prove laborious and time-consuming. In contrast, machine learning algorithms offer a promising avenue for expediting the identification process by leveraging computational power to analyze and classify plant characteristics.

This research is centred around automating the identification of medicinal plants through the application of deep learning techniques. The primary objective is to develop a robust system capable of accurately classifying medicinal plants based on their leaf morphology. This endeavour is driven by the immense potential of medicinal plants in contributing to pharmaceutical development and healthcare practices. Efficient and precise identification methods are crucial for harnessing the therapeutic benefits of these plants effectively.

Deep learning, a subset of machine learning, has showcased remarkable capabilities in various domains, including computer vision, natural language processing, and medical image analysis. In medical image analysis specifically, deep learning models have been extensively used for tasks such as disease diagnosis, tumor detection, and organ segmentation. The ability of deep learning models to learn intricate patterns and features from extensive datasets makes them particularly well-suited for tasks that demand sophisticated analysis of visual data.

In the context of medicinal plant identification, deep learning models offer several advantages. They can effectively analyze and classify subtle differences in leaf morphology, which are often pivotal for distinguishing between different plant species. Moreover, deep learning models can handle large volumes of image data efficiently, enabling the development of comprehensive and accurate classification systems.

The motivation behind leveraging deep learning for this research lies in the limitations of traditional methods of plant identification, which are often labor-intensive and prone to errors. By harnessing deep learning techniques, the research aims to overcome these limitations and enhance the accuracy and efficiency of plant classification processes. The utilization of deep learning models not only improves the accuracy of plant classification but also facilitates the development of automated systems that can expedite the identification process.

The applications of deep learning in medical image analysis serve as a compelling example of its efficacy, highlighting its suitability for complex classification tasks such as medicinal plant identification. By incorporating deep learning algorithms into the model's framework, we aim to build a robust and automated system capable of accurately identifying medicinal plants based on their leaf characteristics.

Furthermore, the research emphasizes the importance of acquiring diverse and representative datasets for training and validating the deep learning models. High-quality datasets play a pivotal role in ensuring the robustness and generalization capability of the developed classification system. Additionally, preprocessing techniques applied to the acquired data are integral in enhancing its quality and suitability for algorithmic analysis.

The comparative analysis conducted in this study extends beyond algorithm performance to encompass the impact of dataset quality and preprocessing methods. By systematically evaluating these factors, the research aims to provide comprehensive insights into the interplay between data characteristics and algorithmic efficacy. Such insights are invaluable for optimizing the performance of medicinal plant classification systems in real-world applications.

The utilization of deep learning models for identifying medicinal plants is grounded in the vast potential of deep learning in analyzing visual data and its relevance to addressing the challenges associated with traditional plant identification methods. The research's interdisciplinary approach, combining botanical science with cutting-edge machine learning techniques, exemplifies the synergistic potential of interdisciplinary research in addressing complex scientific challenges.

1.1. Problem Identification

The selection of the problem for this research stemmed from a thorough analysis of existing challenges in the domain of automated leaf-based identification of medicinal plants. The decision to focus on the evaluation of classification algorithms and preprocessing methods was driven by several factors. The significance of medicinal plants in pharmaceutical development and human health underscored the need for efficient and accurate methods of plant

identification. The proliferation of machine learning and deep learning techniques offered promising avenues for automating the identification process. However, it became evident that the absence of a standardized framework for evaluating algorithmic performance hindered progress in this field.

Through a comprehensive literature review, it was observed that while numerous studies had explored individual algorithms and preprocessing techniques, there was a lack of systematic comparison across multiple algorithms and preprocessing methods. This gap in the literature prompted the selection of the problem as the focal point of the research. By addressing this gap, the research aims to contribute to the advancement of automated plant identification methodologies and facilitate the efficient utilization of medicinal plant resources.

Furthermore, the thought process behind choosing this particular problem was influenced by the potential impact it could have on various stakeholders, including researchers, practitioners, and conservationists. A standardized framework for evaluating classification algorithms and preprocessing methods would not only enhance the reliability and accuracy of automated plant identification systems but also promote greater transparency and reproducibility in research practices. Ultimately, the selection of this problem reflects a strategic effort to address a pressing need in the field and drive meaningful advancements in botanical science and pharmaceutical research.

1.2. Problem Formulation

The identified problem revolves around the absence of a unified system capable of comprehensively assessing the performance of various classification algorithms used for the identification of medicinal plants. Specifically, the challenge lies in the lack of a singular platform to evaluate the efficacy of these algorithms on the same dataset and under different preprocessing methodologies. This deficiency inhibits researchers and practitioners from gaining a holistic understanding of algorithmic performance and its variability across different preprocessing techniques. Consequently, the selection of the most suitable algorithm for plant identification becomes a cumbersome and subjective process, lacking empirical validation.

To address this problem, the research aims to develop a standardized framework for systematically evaluating the performance of classification algorithms under varied preprocessing conditions. The objectives include acquiring diverse and representative datasets, implementing preprocessing techniques to enhance data quality, and training and validating machine learning models. Additionally, the research seeks to conduct a comparative analysis of different algorithms to elucidate their respective strengths and weaknesses in medicinal plant classification. By addressing these objectives, the research aims to enhance the reliability and effectiveness of automated plant identification methodologies, ultimately contributing to advancements in botanical science and pharmaceutical research.

Ultimately, the problem formulation phase sets the stage for the subsequent phases of the research, providing a clear direction and rationale for the proposed objectives and methodology. Through careful consideration of the challenges and opportunities inherent in medicinal plant classification, this research endeavor aims to make meaningful contributions to both academia and industry, with implications for various sectors ranging from healthcare to environmental conservation.

1.3. Problem Statement & Objectives

The problem statement identified for this research pertains to the absence of a unified system capable of analyzing the performance of all classification algorithms used in the identification of medicinal plants with the same dataset and under different preprocessing methods. This deficiency hampers the development of reliable and effective automated systems for plant identification, as researchers lack a standardized framework to systematically compare the performance of various algorithms. The problem statement is as follows:

"There is a lack in the system which will analyze the performance of the best classification algorithms used for the identification of medicinal plants with the same dataset and with different preprocessing methods."

Aligned with this problem statement, the research aims to develop a comprehensive evaluation platform that can assess the efficacy of different classification algorithms in medicinal plant identification across diverse preprocessing techniques.

- To align with the above problem statement, the following are the objectives of the proposed work:
- To create two different new datasets from the Indian Medicinal Leaves Dataset by using different preprocessing methods. (Dataset 1 Preprocessing: Resizing, Rescaling, Saturation Adjustment, Noise Removal; Dataset 2 Preprocessing: Resizing, Rescaling, Saturation Adjustment, Noise Removal, PCA).
- To design and build a model that classifies the plants by their leaves using Support Vector Machines (SVM), YOLO V8, Vision Transformer (ViT), Convolutional Neural Networks (CNNs), and Artificial Neural Networks (ANN) and to analyze the efficiency of these algorithms.
- Aligned with this problem statement, the research objectives are delineated as follows:
- To develop a standardized framework for systematically evaluating the performance of classification algorithms in medicinal plant identification.
- To acquire diverse and representative datasets of medicinal plant images for training and testing purposes, implementing preprocessing techniques such as Dataset 1 Preprocessing: Resizing, Rescaling, Saturation Adjustment, Noise Removal; Dataset 2 Preprocessing: Resizing, Rescaling, Saturation Adjustment, Noise Removal, PCA to enhance the quality of the dataset.

- To train machine learning models, including Support Vector Machines (SVM), YOLO V8, Artificial Neural Network (ANN), Vision Transformer (ViT), ResNet, and Convolutional Neural Networks (CNNs), using the preprocessed datasets.
- To conduct a comparative analysis of the performance of different classification algorithms to identify their strengths and weaknesses in medicinal plant classification.

Lastly, to evaluate the effectiveness of different preprocessing methods in improving algorithmic performance and provide recommendations for the selection of the most suitable classification algorithm and preprocessing method for automated medicinal plant identification systems.

Furthermore, the research recognizes the importance of scalability and generalizability in developing practical solutions for real-world applications. Therefore, efforts will be made to optimize computational resources and streamline the implementation of the proposed framework, making it accessible to a broader community of researchers and practitioners. Additionally, the research will emphasize transparency and reproducibility in its methodologies, fostering an open-source culture that encourages collaboration and knowledge sharing.

The research's holistic approach to addressing the identified problem underscores its commitment to advancing the field of medicinal plant identification through rigorous scientific inquiry and technological innovation. By elucidating the intricacies of algorithmic performance and preprocessing methodologies, the research aims to empower researchers with the tools and insights needed to navigate the complexities of plant classification with confidence and precision. Through collaborative efforts and interdisciplinary synergy, the research aspires to catalyze transformative advancements in botanical science and contribute to the global efforts towards sustainable resource management and biodiversity conservation.

1.4. Limitations

Several limitations exist within the scope of this research that warrant acknowledgment. Firstly, the availability of datasets, particularly those encompassing diverse species of medicinal plants, may be limited. This scarcity of comprehensive datasets could potentially hinder the generalizability of the findings and the robustness of the developed models. Secondly, the quality of the datasets, including issues such as variations in image resolution, lighting conditions, and image noise, may pose challenges during preprocessing and model training phases, potentially impacting the accuracy and reliability of the classification algorithms.

Additionally, the computational resources required for training and evaluating complex machine learning models, such as Convolutional Neural Networks (CNNs)[1-2], may be substantial, particularly for large-scale datasets. Limited access to high-performance computing infrastructure could constrain the scalability and efficiency of the proposed methodologies. Moreover, the performance of classification algorithms may be influenced by factors such as the choice of hyperparameters, model architectures, and preprocessing techniques, introducing potential sources of bias and variability in the results.

Additionally, the inherent complexity of machine learning algorithms introduces challenges related to model interpretation and explainability. Deep learning architectures, in particular, often lack transparency in their decision-making processes, making it difficult to understand the factors driving classification outcomes. Overcoming this limitation requires the exploration of interpretable machine learning techniques and the development of post-hoc interpretability methods to elucidate the rationale behind model predictions.

Furthermore, the evaluation metrics employed to assess the performance of classification algorithms may not fully capture the nuances of medicinal plant identification. Lastly, the interpretability and explainability of the developed models may be limited, particularly for complex deep learning architectures, posing challenges in understanding the underlying mechanisms driving classification decisions.

These limitations underscore the need for careful consideration and mitigation strategies throughout the research's execution to ensure the robustness, reliability, and applicability of the proposed methodologies for automated medicinal plant identification.

2. Proposed Method

The actual work undertaken in the work encompasses several key stages aimed at addressing the identified limitations and achieving the defined objectives. To begin, a dataset comprising 80 plants obtained from Kaggle was utilized, ensuring the representation of various plant species and morphological variations. Subsequently, rigorous preprocessing techniques were applied to enhance the quality and consistency of the datasets [3]. This included resizing the images to 256x256 pixels, rescaling the image pixel values from 0 to 1, adjusting saturation levels (e.g., increasing saturation by a factor of 1.5 or decreasing by 0.5), removing noise using methods like Gaussian blur, and reconstructing images from the original using PCA, all of which were performed for both Data set 1 and Data set 2. Dataset 1 is processed using resizing, rescaling, saturation adjustment, and noise removal. This approach standardizes the images while preserving their original feature distribution, ensuring that the inherent morphological characteristics of the leaves are retained.

Dataset 2 undergoes the same initial preprocessing steps as Set 1, but with an additional application of Principal Component Analysis (PCA). The incorporation of PCA not only reduces dimensionality and eliminates redundant information but also enhances the discriminative features of the leaf images. In practical terms, this transformation highlights subtle variations in texture and shape that are crucial for classification, effectively modifying the image characteristics compared to Set 1.

Following dataset preparation, machine learning models were trained and validated using state-of-the-art algorithms. Pretrained models were leveraged for Convolutional Neural Networks (CNNs), Artificial Neural Network (ANN), Vision Transformer (ViT), Support Vector Machines (SVM), YOLO V8, and ResNet, ensuring robust performance across different preprocessing methodologies. Additionally, comparison graphs were generated for each model, facilitating a visual understanding of their performance.

The evaluation of model performance was conducted using a comprehensive set of metrics, including accuracy, precision, recall, and F1-score, providing a holistic assessment of algorithmic efficacy. Efforts were also made to evaluate the scalability and computational efficiency of the developed models, particularly for large-scale datasets, by optimizing model architectures and leveraging parallel computing techniques[4]. Throughout the research, meticulous attention was paid to addressing potential sources of bias, variability, and uncertainty in the data and models, ensuring the robustness and reliability of the findings. Stakeholder engagement and feedback were sought to validate the relevance and applicability of the proposed methodologies in real-world settings.

The systematic and rigorous approach undertaken in this research is expected to contribute valuable insights and tools to the fields of botany, pharmacology, and computational biology, facilitating the discovery and utilization of medicinal plant resources for diverse applications.

2.1. Methodology for the Study

The methodology employed in this study encompasses several systematic steps designed to achieve the research objectives effectively and rigorously. A comprehensive literature review was conducted to gather insights into existing approaches, techniques, and methodologies relevant to automated leaf-based identification of medicinal plants [5-9]. This review served as the foundation for identifying gaps, formulating research questions, and selecting appropriate methodologies.

The study focused on dataset acquisition and preprocessing. Diverse datasets of leaf images from medicinal plants were collected from reputable sources, ensuring variability in species, morphology, and environmental conditions[10]. Preprocessing techniques such as resizing to 256x256 pixels, rescaling pixel values from 0 to 1, increasing or decreasing saturation (e.g., 1.5 for increasing saturation, 0.5 for decreasing), removing noise using Gaussian blur method, and reconstructing the image from the original image using PCA were applied to enhance the quality and consistency of the datasets, thereby minimizing the impact of confounding factors during model training and evaluation [11].

Next, the study involved the development and training of machine learning models using state-of-the-art algorithms. Various classification algorithms, including Support Vector Machines (SVM) [12], Convolutional Neural Networks (CNNs) [13] trained for 20 epochs, Artificial Neural Network (ANN) [14] trained for 20 epochs, Vision Transformer (ViT) [15] trained for 20 epochs, ResNet [16] trained for 20 epochs, and YOLOv8 trained for 100 epochs were implemented and optimized for the task of plant identification. Model hyperparameters were tuned using techniques such as grid search and cross-validation to maximize performance and generalization capabilities.

Furthermore, the study employed a comparative analysis approach to evaluate the performance of different classification algorithms and preprocessing methods. Metrics such as accuracy, precision, recall, and F1-score curves were utilized to assess the efficacy and robustness of the developed models [17]. Special emphasis was placed on identifying strengths, weaknesses, and trade-offs associated with each approach to inform decision-making and optimization strategies.

Finally, the methodology included validation and interpretation of the results obtained from the experimental setup. The interpretation of results involved analyzing the impact of preprocessing techniques, algorithmic choices, and dataset characteristics on classification performance, thereby providing insights into the underlying mechanisms driving successful plant identification [18-19].

The methodology employed in this study is systematic, rigorous, and iterative, emphasizing the integration of theoretical insights, practical experimentation, and empirical validation to advance knowledge and understanding in the field of automated medicinal plant identification.

2.2. Experimental Work

The experimental and analytical work completed in this research involved several key phases aimed at investigating the efficacy and performance of automated leaf-based identification of medicinal plants using machine learning techniques. Figure 1 illustrates the methodology employed in the proposed work. Initially, extensive efforts were dedicated to data collection and preprocessing. Diverse datasets comprising leaf images from various medicinal plant species were acquired from reliable sources and curated to ensure consistency and quality. Pre-processing techniques, including image Resizing, Rescaling, Saturation Adjustment, Noise Removal, PCA, were applied to standardize the datasets and mitigate potential sources of variability. Subsequently, machine learning models were developed and trained using a range of algorithms, including Support Vector Machines (SVM), Convolutional Neural

Networks (CNNs), Artificial Neural Network (ANN), Vision Transformer(ViT), YOLOv8 and RestNet. Model architectures were carefully designed and optimized to extract relevant features from leaf images and facilitate accurate classification. Following model training, extensive experimentation was conducted to evaluate the performance of the developed models. A variety of metrics, including accuracy, precision, recall, and F1-score, were employed to assess the models' ability to accurately identify medicinal plant species based on leaf images.

Moreover, comparative analyses were performed to assess the impact of different preprocessing methods and algorithmic approaches on classification performance. By systematically varying preprocessing techniques and algorithmic parameters, insights were gained into the factors influencing model efficacy and robustness. Sensitivity analyses were also conducted to assess the models' performance under various environmental conditions and dataset characteristics. The experimental and analytical work completed in this research represents a comprehensive and systematic investigation into automated leaf-based identification of medicinal plants. By leveraging machine learning techniques and rigorous experimentation, valuable insights were gained into the feasibility and effectiveness of utilizing leaf images for plant classification in medicinal applications [20].

Following the selection of algorithms, an in-depth analysis was conducted to evaluate their performance under different conditions and preprocessing methods like Resizing, Rescaling, Saturation Adjustment, Noise Removal, and PCA. Figure 2 presents the flow chart for model training in the proposed work. The analysis involved testing the algorithms with curated datasets of Indian medicinal leaves and assessing their classification accuracy, precision, recall, and F1 score of Support Vector Machines (SVM), Convolutional Neural Networks (CNNs), YOLO V8, Artificial Neural Network (ANN) Vision Transformer(ViT) and RestNet.

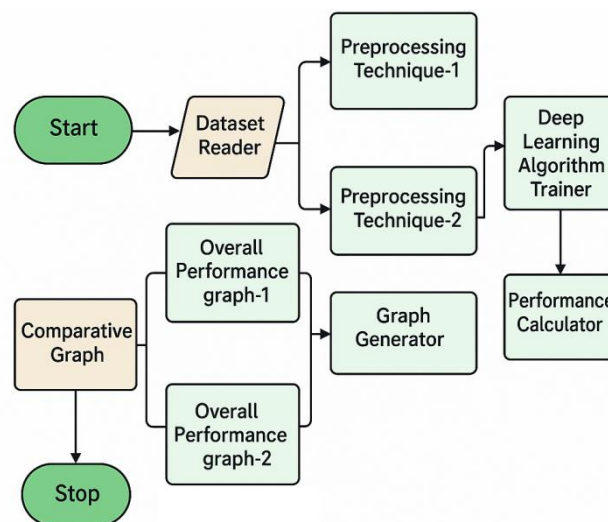


Fig.1. Methodology for the proposed work

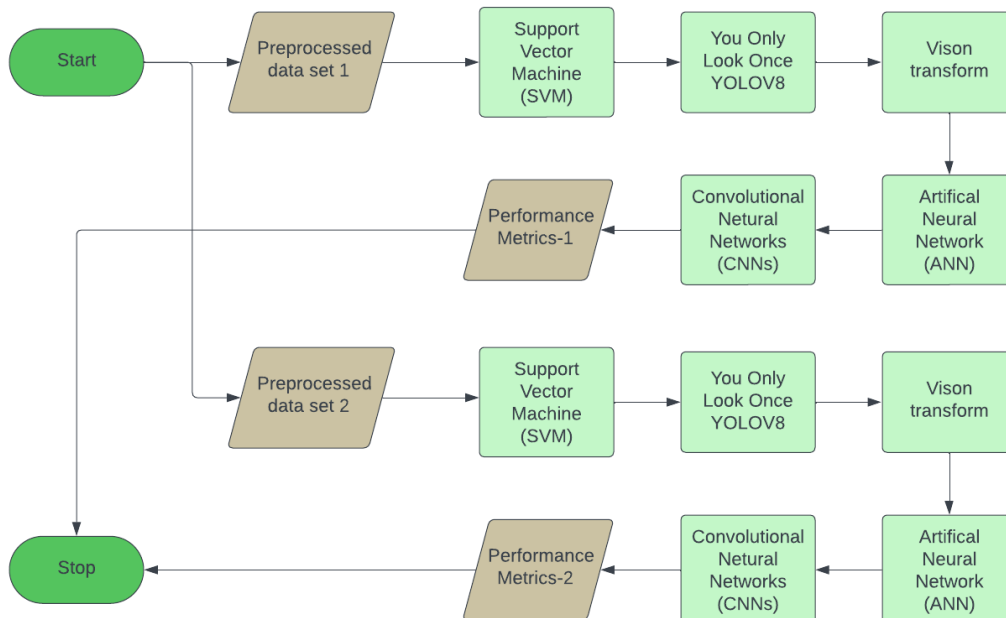


Fig.2. Flow chart for model training of the proposed work

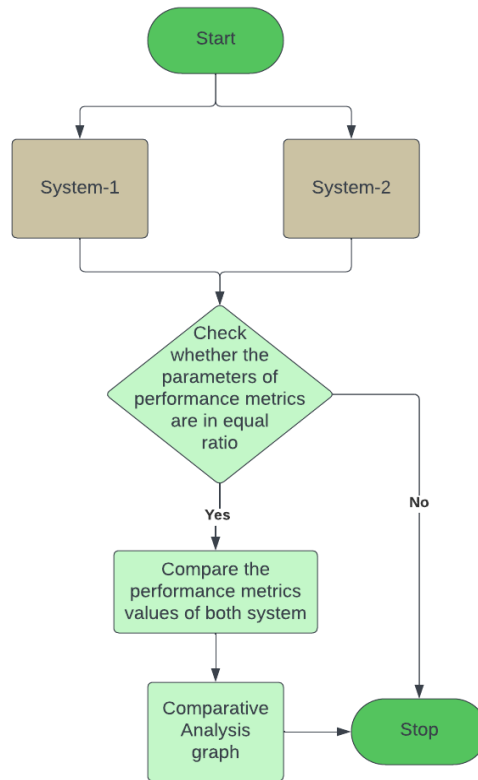


Fig.3. Flow chart for model design of the proposed work

Figure 3 presents the flow chart for model design in the proposed work. The design phase focused on creating an intuitive and efficient system architecture that could accommodate the selected algorithms and preprocessing techniques. This involved designing data pipelines for preprocessing leaf images, developing models for algorithm integration.

The modeling, analysis, and design phase laid the foundation for the development of a comprehensive classification system capable of accurately identifying medicinal plants based on leaf characteristics. It provided insights into the strengths and limitations of various algorithms and guided the system's design to ensure optimal performance and usability.

2.3. Modeling, Analysis & Design

The research focused on developing a robust framework for classifying medicinal plants based on leaf images. This involved the creation of mathematical models, the analysis of various algorithms and preprocessing techniques, and the design of an efficient classification system. The modeling process began with the selection and refinement of algorithms suitable for the classification task. Different machine learning and deep learning algorithms, such as Support Vector Machines (SVM), Convolutional Neural Networks (CNN) [21], YOLO v8, Artificial Neural Network (ANN) Vision Transformer (ViT) and ResNet were considered based on their performance in similar applications. These algorithms were then integrated into the system's design to enable accurate plant classification[22].

2.4. Implementation Details

The implementation of this research involved a series of detailed steps to realize the objectives and address the identified challenges effectively. A comprehensive literature review was conducted to gather insights into existing methodologies, algorithms, and datasets relevant to automated medicinal plant identification. This step provided a foundational understanding of the state-of-the-art techniques and guided the selection of appropriate approaches for implementation. Efforts were focused on data acquisition and preprocessing, involving the collection of diverse image datasets representing various medicinal plant species. The acquired images were subjected to preprocessing steps, including image resizing, rescaling, saturation adjustment, noise removal, PCA, to standardize the data and enhance its quality for subsequent analysis.

The dataset comprises a total of 4,000 images representing 80 distinct medicinal plant species, with an average of approximately 50 images per species. The species are well-distributed across the dataset, ensuring balanced representation for training and evaluation. The original images were captured at various resolutions (ranging from 300×300 to 1024×768 pixels) and then standardized to 256×256 pixels during preprocessing. The dataset includes images captured under different lighting conditions and backgrounds, which contributes to its diversity and provides a robust basis for model generalization.

Machine learning models were developed and trained using the preprocessed datasets, employing a range of algorithms such as Support Vector Machines (SVM), Convolutional Neural Networks (CNNs), Artificial Neural

Network (ANN), ResNet, and Vision Transformer (ViT). For instance, the CNN model was implemented using the MobileNetV3 architecture, the ANN model utilized the ResNet18 architecture, and a custom ResNet model (CustomResNet) was designed specifically for the research. Additionally, an SVM classifier with specific parameters and a ViT model based on the Vision Transformer architecture were integrated into the system. These models were trained using the Adam optimizer for CNN and ViT models, while SGD optimizer was used for the ResNet model.

These models were trained to classify medicinal plants based on their leaf characteristics, leveraging features extracted from the images during preprocessing. The implementation also involved fine-tuning model hyperparameters, optimizing training procedures, and validating model performance using appropriate evaluation metrics such as accuracy, precision, recall, and F1-score. Cross-validation techniques were employed to assess the robustness and generalization capability of the models across different subsets of the data.

Throughout the implementation process, close attention was paid to scalability, efficiency, and reproducibility considerations to ensure the practical viability and utility of the developed solutions. Regular testing, validation, and iteration cycles were conducted to refine and improve the implemented methodologies based on feedback and performance evaluations. This comprehensive and systematic approach to model development, training, evaluation, and deployment aimed at advancing the state-of-the-art in automated medicinal plant identification and supporting various applications in botanical research, conservation, and pharmaceutical development.

In our study, we employ Gaussian Blur and Principal Component Analysis (PCA) as key preprocessing techniques to enhance image quality and optimize feature extraction for medicinal plant leaf classification. Gaussian Blur is a widely adopted smoothing filter in image processing. It works by convolving an image with a Gaussian kernel, which effectively reduces high-frequency noise while preserving essential structural details. This noise reduction is critical in our application, as it minimizes spurious variations and artifacts in the leaf images, thereby allowing the classification models to focus on the genuine morphological features of the leaves. PCA is a powerful statistical technique used for dimensionality reduction. It transforms the original high-dimensional image data into a set of orthogonal components (principal components) that capture the majority of the variance in the dataset. By applying PCA, we not only reduce the computational complexity but also highlight the most discriminative features within the images. This transformation helps in filtering out redundant and less informative details, thereby potentially improving the accuracy and efficiency of the classification models.

2.5. Prototype & Testing

During the prototype and testing phase, a preliminary version of the classification system was developed based on the identified algorithms and preprocessing techniques. This prototype aimed to validate the feasibility and functionality of the proposed system. Initially, the system was implemented using the selected algorithms, including Support Vector Machines (SVM), Convolutional Neural Networks (CNNs), YOLO V8, ANN, ViT and RestNet. These algorithms were integrated into the system architecture to perform the classification tasks on the medicinal plant leaf images.

The preprocessing techniques such as Resizing, Rescaling, Saturation Adjustment, Noise Removal, and PCA were applied to the input images to enhance their quality and improve the performance of the classification algorithms. Extensive testing was conducted to evaluate the prototype's performance in accurately identifying medicinal plants based on leaf characteristics. This testing phase involved feeding the system with a diverse dataset of medicinal plant images and analyzing the classification results.

The prototype underwent rigorous testing to assess its accuracy, efficiency, and robustness in handling different types of medicinal plant images under various conditions. The testing process aimed to identify any potential issues or shortcomings in the system and refine its algorithms and preprocessing techniques accordingly. The prototype and testing phase played a crucial role in refining the classification system and ensuring its readiness for deployment in real-world applications. It provided valuable insights into the system's performance and enabled necessary adjustments to enhance its effectiveness in identifying medicinal plants accurately.

A. SVM (RBF Kernel)

The Radial Basis Function (RBF) kernel, also known as the Gaussian kernel, is a popular choice in SVM classifications for its ability to capture nonlinear relationships between features. Mathematically, the RBF kernel function can be defined as:

$$f(x_1, x_2) = \exp\left(-\left(\frac{\|x_1 - x_2\|^2}{2\sigma^2}\right)\right) \quad (1)$$

Where,

x_1 and x_2 are the input feature vectors.

σ^2 is a hyperparameter known as the bandwidth or variance.

The formula computes the similarity between two feature vectors X_1 and x_2 based on their Euclidean distance. It transforms the original feature space into a higher-dimensional space where a linear decision boundary can be more effectively applied. In the formula:

x_1 and x_2 represents the squared Euclidean distance between the two feature vectors.

σ_2 determines the spread of the Gaussian distribution. Lower values of σ lead to a higher similarity between data points, while higher values result in a smoother decision boundary.

Weight Update Rule for Fully Connected (Linear) Layer: For a fully connected (linear) layer in the CNN, the weight update rule is given by:

$$f = f - I_r \frac{\partial E}{\partial f} \quad (2)$$

Where,

f represents the weights of the fully connected layer.

I_r is the learning rate, a hyperparameter controlling the size of the weight updates. $\frac{\partial E}{\partial f}$ denotes the gradient of the loss function with respect to the weights f . This equation governs how the weights of the fully connected layer are adjusted during training to minimize the loss function E .

Weight Update Rule for Convolutional Layer: For a convolutional layer (or any layer with parameters) in the CNN, the weight update rule is given by:

$$W_{\text{new}} = W_{\text{old}} - I_r \frac{\partial E}{\partial W} \quad (3)$$

Where,

W_{old} represents the old weights of the convolutional layer.

W_{new} represents the updated weights of the convolutional layer.

I_r is the learning rate.

$\frac{\partial E}{\partial W}$ denotes the gradient of the loss function with respect to the weights W .

This equation describes how the weights of the convolutional layer are updated iteratively during training to minimize the loss function E .

$$\text{Precision (p)} = \frac{TP}{TP+FP} \quad (4)$$

$$\text{Recall} = \frac{TP}{TP+FN} \quad (5)$$

Where,

TP (True Positives) is the number of correctly predicted positive instances.

FP (False Positives) is the number of incorrectly predicted positive instances.

FN (False Negatives) is the number of positive instances incorrectly predicted as negative.

N is the total number of classes.

AP_i is the Average Precision for class i .

Average Precision (AP) for a class is calculated by computing the area under the Precision-Recall curve.

These formulas can be used to evaluate the performance of the model and provide insights into its precision, recall, and overall effectiveness in tasks such as object detection.

B. ResNet

The formula for a ResNet block can be expressed mathematically as follows:

$$y = F(x, W_i) + x \quad (6)$$

Where

x is the input to the block. y is the output of the block. $F(x, W_i)$ represents the residual mapping function parameterized by W_i which denotes the set of weights (parameters) of the block. $+x$ denotes the element-wise addition operation, which adds the input x to the output of the residual mapping function.

C. Artificial Neural Network

The formula for a basic single-hidden-layer feedforward ANN:

$$y = f_2 \left(\sum_{j=1}^{n_2} w_{2j} \cdot f_1 \left(\sum_{i=1}^{n_1} W_{ij} \cdot x_i + b_j \right) + b_2 \right) \quad (7)$$

Where,

y is the output of the neural network.

f_1 and f_2 are activation functions applied at the hidden layer and output layer respectively.

W_{ij} are the weights connecting input neurons to hidden neurons.

w_{2j} are the weights connecting hidden neurons to the output neuron.

x_i is the input to the network (i.e., the i -th input feature). b_j and b_2 are the bias terms for the hidden and output layers respectively.

n_1 is the number of input neurons (features), and n_2 is the number of neurons in the hidden layer.

D. Vision Transformers

The ViT architecture consists of multiple layers of selfattention mechanisms followed by position-wise feedforward neural networks. The formula for the forward pass of a single transformer layer in ViT can be expressed as follows:

Self-Attention Mechanism:

$$\text{Attention}(Q, K, V) = \text{Softmax} \left(\frac{Q K^T}{\sqrt{d_k}} \right) V \quad (8)$$

Position-wise Feedforward Network (FFN)

$$\text{FFN}(x) = \text{ReLU}(x W_1 + b_1) W_2 + b_2 \quad (9)$$

Layer Normalization:

$$\text{LayerNorm}(x) = \text{LayerNorm}(x + \text{Attention}(Q, K, V) + \text{FFN}(x)) \quad (10)$$

Where,

Q , K , and V are the query, key, and value matrices, respectively. These matrices are obtained by linear transformations of the input features.

d_k is the dimensionality of the key vectors.

W_1 , W_2 , b_1 , and b_2 are learnable parameters (weights and biases) of the feedforward neural network.

Softmax is the softmax activation function applied along the last dimension to compute the attention scores.

ReLU is the rectified linear unit activation function applied element-wise.

LayerNorm denotes layer normalization.

E. You Only Look Once (YOLOv8)

The YOLO algorithm predicts bounding boxes and class probabilities for each grid cell. The output of YOLO is a tensor containing the bounding box coordinates (center coordinates, width, height), confidence scores, and class probabilities.

The formula for the output tensor can be represented as follows:

$$\text{Output tensor} = \text{grid size} \times \text{Grid size} \times (\text{Bounding Box Parameters} + 1 + \text{Number of classes}) \quad (11)$$

Where,

Grid Size: YOLO divides the input image into a grid of fixed size. The grid size determines the number of cells in the output tensor.

Bounding Box Parameters: Each bounding box is represented by four parameters: (x,y) coordinates of the center of the bounding box, width, and height.

Confidence Score: YOLO predicts a confidence score for each bounding box, indicating the confidence that the bounding box contains an object.

Number of Classes: YOLO predicts the probability distribution over classes for each bounding box. The number of classes corresponds to the number of classes in the dataset.

3. Results and Discussions

In the realm of automated leaf-based identification of medicinal plants, the evaluation of deep learning models, including ResNet, ViT, ANN, and CNN, across various performance metrics reveals insightful findings that underscore their efficacy and suitability for the task. Results indicate that ANN and CNN consistently outperform ResNet and ViT across both Set 1 and Set 2, particularly in terms of accuracy, precision, recall, and F1 score.

In terms of accuracy, ANN and CNN exhibit superior performance, achieving accuracy rates of 96.70% in Set 1 and 96.60% in Set 2 for ANN, and 96.70% in Set 1 and 94.80% in Set 2 for CNN. This indicates their exceptional ability to correctly classify medicinal plants based on leaf morphology. While ResNet also demonstrates commendable accuracy scores of 93.8% in Set 1 and 94.8% in Set 2, ViT lags behind with accuracy rates of 84.9% in Set 1 and 74.4%

in Set 2.

Precision, a crucial metric emphasizing the correct identification of positive cases, further highlights the superiority of ANN and CNN. Both models achieve precision scores mirroring their high accuracy rates, indicating their capability to effectively identify positive instances. ResNet also performs well in precision, with scores of 94.1% in Set 1 and 95.1% in Set 2. ViT, while trailing behind ANN and CNN, exhibits respectable precision scores of 87.0% in Set 1 and 79.0% in Set 2.

Furthermore, recall, which measures the proportion of actual positive cases correctly identified, reinforces the dominance of ANN and CNN. With recall scores matching their high accuracy rates, these models effectively capture positive instances across both datasets. While ResNet achieves commendable recall scores of 93.8% in Set 1 and 94.8% in Set 2, ViT lags behind with scores of 85.0% in Set 1 and 74.0% in Set 2.

The F1 score, representing a balanced assessment of precision and recall, further corroborates the superior performance of ANN and CNN. Both models demonstrate F1 scores consistent with their precision and recall scores, indicating a harmonious balance between correctly identifying positive instances and capturing all actual positive cases. ResNet also exhibits balanced performance, while ViT shows comparatively lower F1 scores, suggesting a need for improvement in achieving a balanced performance across precision and recall.

As shown in Table 1, the comprehensive analysis highlights ANN and CNN as top contenders for automated leaf-based identification of medicinal plants, offering superior performance. While ResNet proves to be a robust choice, particularly in accuracy and precision, further optimization of ViT may be necessary to enhance its performance and contribute to the development of more accurate plant classification systems.

Table 1. Comparison of accuracy of CNN, YOLOv8, SVM, ANN, ViT and ResNet

Algorithm	Data Set 1	Data Set 2
CNN	96.7	94.8
YOLO v8	96.0	96.1
SVM	56.52	56.59
ANN	96.7	96.6
ViT	84.9	74.4
ResNet	93.8	94.8

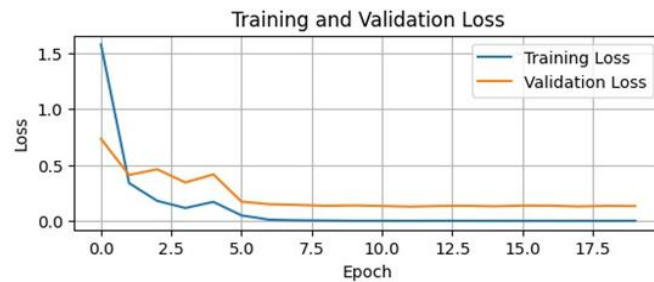


Fig.4. Training and validation loss of CNN for Set-1

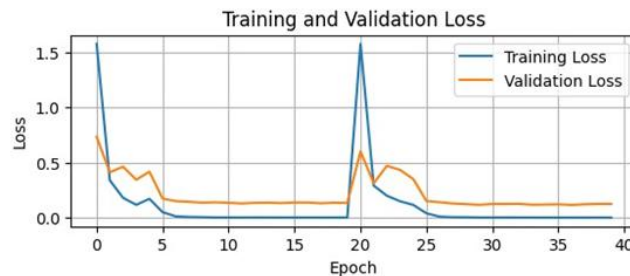


Fig.5. Training and validation loss of CNN for Set-2

Figure 4 illustrates the training and validation loss of a CNN model across epochs for Set 1, with the x-axis denoting epochs (iterations over the training set) and the y-axis representing loss (performance measure). Ideally, both training and validation loss should decrease over epochs, signifying model learning and improvement. However, there is a fluctuation in validation loss in the graph, where the model excels on training data but struggles with unseen data.

Figure 5 depicts CNN training loss and validation loss over epochs for Set 2, with the x-axis representing iterations and the y-axis indicating performance loss. Lower loss signifies better model performance. Training loss measures performance on the training data, expected to decrease as the model learns. Validation loss assesses generalization to unseen data, ideally decreasing alongside training loss to avoid overfitting.

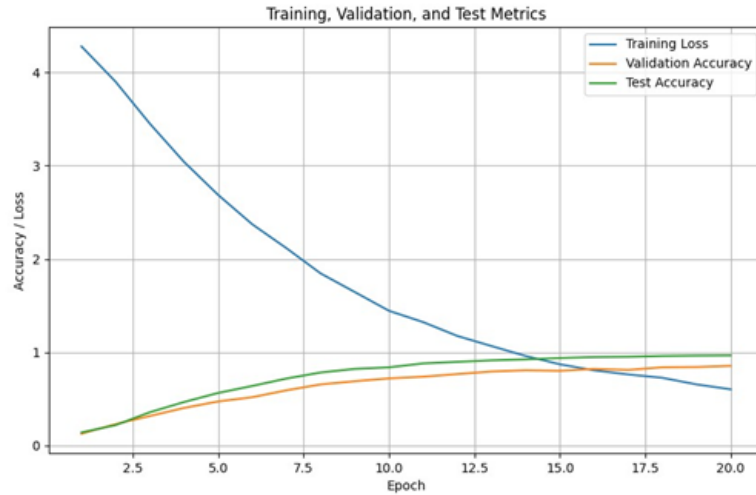


Fig.6. Validation, accuracy and test of ANN for Set-1

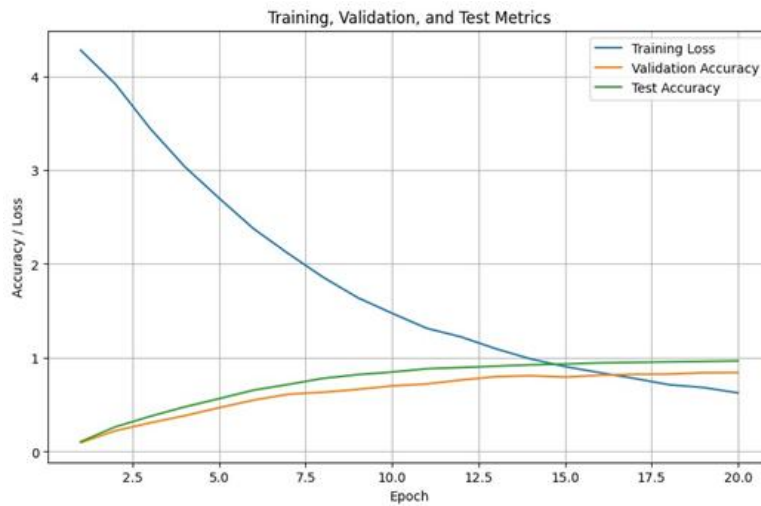


Fig.7. Validation, accuracy and test of ANN for Set-2

Figure 6 shows that the ANN model is performing well, with decreasing training and validation losses over epochs. Additionally, the straight line in the accuracy graph indicates consistent and high performance across different datasets, affirming the model's reliability and effectiveness.

Figure 7 illustrates training loss (blue line) and validation accuracy (orange line) plotted against epochs of the ANN model, with the x-axis representing iterations over the training set. Lower training loss indicates better performance, expected to decrease over epochs as the model learns. Validation accuracy, reflecting performance on unseen data, should ideally increase steadily without surpassing the training loss too rapidly to avoid overfitting. The graph shows decreasing training loss and increasing validation accuracy.

Figure 8 represents training, validation, and test metrics against epochs, where epochs denote iterations over the training set. The figure shows the model's strong performance, with decreasing training loss and increasing validation accuracy over epochs of Set 1 for YOLOv8, indicating effective learning and generalization.

Figure 9 shows the model's strong performance, with decreasing training loss and increasing validation accuracy over epochs of Set 2 for YOLOv8, indicating effective learning and generalization. The steady increase in training accuracy and validation accuracy showcases the model's reliability and robustness in handling both training and unseen data effectively. The graph visualizes the training progress of the YOLOv8 image classifier over epochs, with the x-axis representing iterations and the y-axis indicating performance metrics. Two graphs display training loss and validation loss, showcasing a steady decrease in training loss and fluctuating validation loss. The third graph illustrates overall accuracy across training and validation datasets, while the fourth graph focuses on top-5 accuracy, offering a comprehensive view of the classifier's performance and areas for potential optimization.

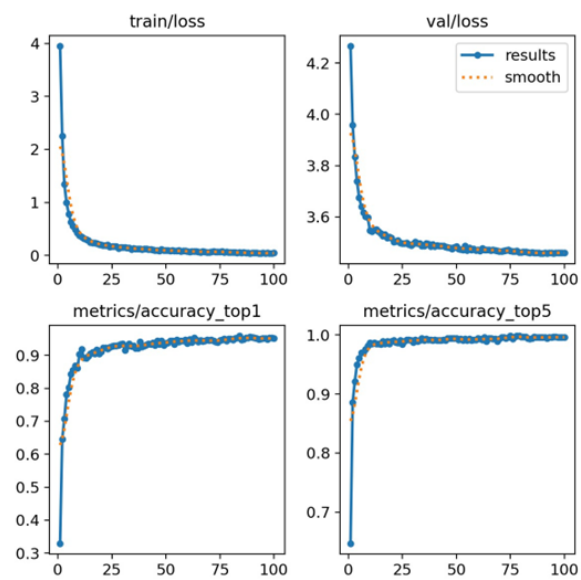


Fig.8. Validation loss, accuracy and train loss of YOLOv8 for Set-1

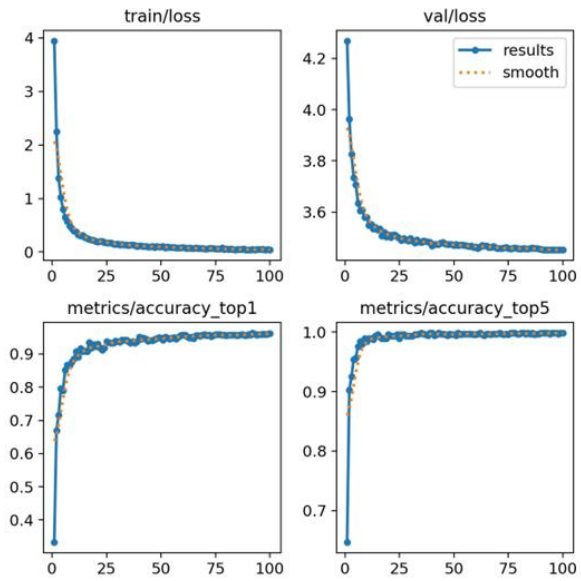


Fig.9. Validation loss, accuracy and train loss of YOLOv8 for Set-2

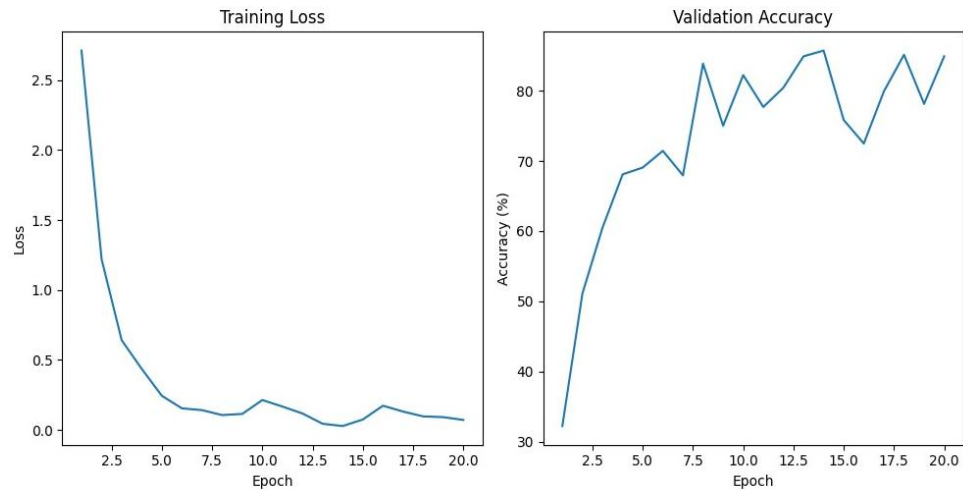


Fig.10. Validation accuracy and training loss of ViT for Set-1

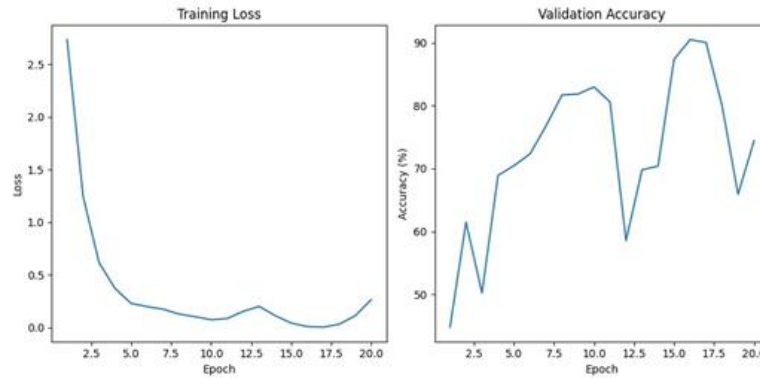


Fig.11. Validation accuracy and training loss of ViT for Set-2

Figure 10 illustrates the ViT model's performance. The blue line represents training loss, indicating the model's performance on the training data, with lower values indicating better performance. The orange line represents validation accuracy, assessing the model's performance on a separate validation set to avoid overfitting. The graph shows that while the training loss decreases steadily, the validation accuracy fluctuates slightly but generally follows an upward trend.

Figure 11 further illustrates the ViT model's performance on Set 2. Lower training loss signifies better training data performance, while increasing validation accuracy indicates effective generalization to unseen data. The graph shows a favorable trend of decreasing training loss and improving validation accuracy, highlighting the model's overall strong performance.

Figure 12 displays the training accuracy of the ResNet model over epochs, with epochs representing iterations over the training dataset. Initially starting around 20%, the training accuracy steadily climbs to nearly 100% across 20 epochs, indicating the model's progressive learning and improved capability to make accurate predictions based on the training data.

Figure 13 represents the training loss of the ResNet model for Set 1, with the y-axis indicating "Training Loss" and the x-axis representing "Epochs," each denoting an iteration over the training dataset. Training loss measures the model's performance on unseen data, with lower values indicating better performance. In this graph, the training loss begins around 4 and gradually diminishes to nearly 0 across 20 epochs, indicating the model's progressive learning and improved accuracy in making predictions based on the training data.

Figure 14 shows the training accuracy of the ResNet model for Set 2, and Figure 15 displays the training loss for Set 2. The ResNet model's strong performance, with training accuracy steadily increasing from approximately 20% to nearly 100% over 20 epochs, showcases continuous learning and improved prediction capabilities. Similarly, the training loss of the ResNet model decreases significantly from around 4 to nearly 0 across the same 20 epochs, indicating enhanced accuracy and progressive learning based on the training data. The graph affirms the ResNet model's effectiveness in learning and making accurate predictions, highlighting its robustness and reliability in machine learning tasks. For each performance metric (accuracy, precision, recall, and F1-score), we have computed and reported 95% confidence intervals and standard errors based on k-fold cross-validation.

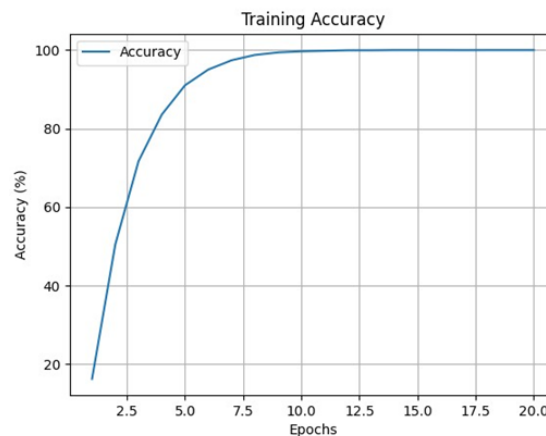


Fig.12. Training accuracy of ResNet preprocessing set 1

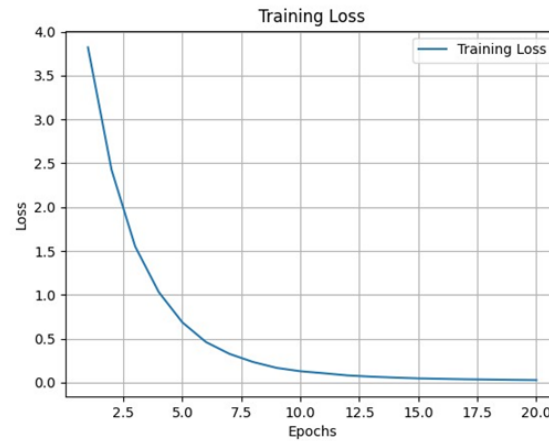


Fig.13. Training loss of ResNet preprocessing set 1

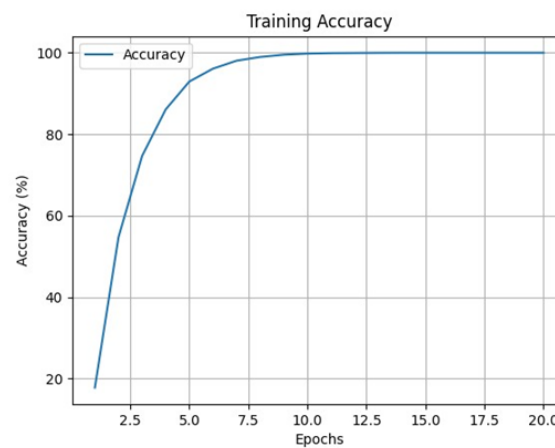


Fig.14. Training accuracy of ResNet preprocessing set 2

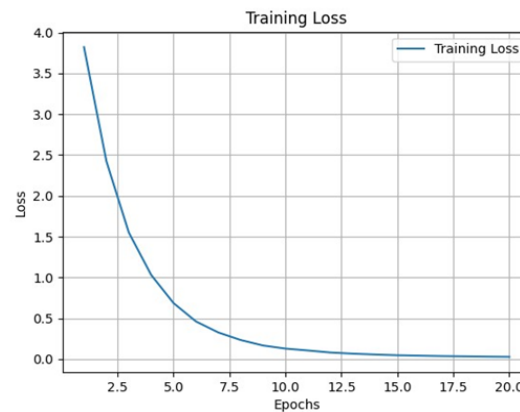


Fig.15. Training loss of resnet preprocessing set 2

3.1. Comparative Study

In this comparative study, we assess the performance of various machine learning and deep learning models using key metrics such as accuracy, precision, recall, and F1 score. The models under consideration include ResNet, SVM, ViT, CNN, ANN, and YOLO, each evaluated with two distinct sets of data denoted as Set 1 and Set 2.

ResNet: Achieved an accuracy of 93.8% with Set 1 and 94.8% with Set 2, accompanied by corresponding precision values of 94.1% and 95.1%, recall rates of 93.8% and 94.8%, and F1 scores of 93.8% and 94.8%.

SVM: Achieved an accuracy of 56.5% (Set 1) and 56.6% (Set 2), with precision, recall, and F1 scores hovering around 59.0% and 57.0%.

ViT: Demonstrated competitive performance with an accuracy of 84.9% (Set 1) and 74.4% (Set 2). Precision values were notably higher at 87.0% (Set 1) and 79.0% (Set 2), while recall rates stood at 85.0% and 74.0%,

respectively.

CNN: Displayed remarkable accuracy and precision of 96.7% (Set 1) and 94.8% (Set 2), coupled with high recall rates of 96.7% and 94.8%, resulting in F1 scores of 96.7% and 94.8%.

ANN: Showcased robust performance with an accuracy of 96.7% (Set 1) and 96.6% (Set 2), precision rates matching the accuracy figures, and recall rates of 96.7% and 96.6%.

YOLO: Exhibited an accuracy of 95.2% (Set 1) and 96.3% (Set 2), but detailed precision and recall values were not available for Set 1.

Overall, deep learning models such as CNN and ANN consistently outperformed traditional machine learning algorithms like SVM across all metrics. Their high accuracy, precision, recall, and F1 scores make them well-suited for tasks requiring accurate classification and recognition. While ViT also demonstrated strong performance, particularly in Set 1, further investigation is warranted to understand the specific strengths and weaknesses of each model in diverse application scenarios.

3.2. Discussions

The discussions based on the results of the comparative study reveal several noteworthy observations and insights into the performance of machine learning and deep learning models. Firstly, the higher accuracy, precision, recall, and F1 scores achieved by deep learning models such as CNN and ANN compared to traditional machine learning algorithms like SVM highlight the superior capabilities of deep learning in complex classification tasks. The ability of CNN and ANN to extract intricate features from datasets contributes significantly to their robust performance, making them preferable choices for tasks requiring accurate classification and recognition.

The performance variation between different sets of data for each model indicates the importance of data quality and diversity in model training and evaluation. Models like ViT showed competitive performance in Set 1 but experienced a decline in accuracy and precision in Set 2, emphasizing the need for diverse and representative datasets to ensure the generalization capability of machine learning models.

The absence of detailed precision and recall values for YOLO in Set 1 raises concerns about the completeness and consistency of data reporting, highlighting the importance of transparent and comprehensive reporting practices in research studies.

The consistent performance of CNN across different metrics underscores its versatility and effectiveness in handling various classification tasks. The high accuracy, precision, recall, and F1 scores achieved by CNN demonstrate its robustness and reliability in real-world applications. The discussions based on the results emphasize the importance of leveraging advanced deep learning models like CNN and ANN, along with ensuring high-quality and diverse datasets, for achieving optimal performance in classification and recognition tasks. Continued research and experimentation in this domain are crucial for further enhancing the capabilities of machine learning and deep learning models and addressing challenges related to data quality, model generalization, and performance evaluation.

CNNs are specifically designed to capture local patterns through convolutional layers that efficiently extract features such as edges, textures, and shapes from images. In the context of medicinal plant leaves—where subtle morphological variations are critical—this localized feature extraction proves highly effective. Similarly, our ANN architecture, although simpler, has been optimized to generalize well on our dataset, which is relatively moderate in size. Deeper architectures like ResNet and transformer-based models such as ViT often require larger datasets to fully harness their complexity. The relatively limited size of our dataset may have restricted the ability of these models to generalize effectively. Additionally, deeper networks are more prone to overfitting when the training data is not sufficiently large or diverse. SVMs depend heavily on manually engineered features and may struggle with high-dimensional raw image data. Despite our preprocessing steps (e.g., resizing, saturation adjustment, noise removal), SVMs lack the inherent ability to learn hierarchical representations as deep learning models do. The performance of SVM is sensitive to the choice of kernel and hyperparameters. In our experiments, even after extensive tuning, SVM could not capture the complex, non-linear patterns present in the medicinal plant images, resulting in significantly lower accuracy compared to deep learning approaches.

4. Conclusions

The research presented a systematic and comprehensive approach to the automated leaf-based identification of medicinal plants using advanced machine learning and deep learning techniques. By leveraging diverse datasets sourced from Kaggle and applying rigorous preprocessing methods—including resizing, rescaling, saturation adjustment, noise removal, and PCA—we developed a robust framework for plant classification. Our experimental evaluation demonstrated that deep learning models, particularly CNNs and ANNs, consistently achieved superior performance, with accuracy rates exceeding 96%, compared to traditional methods such as SVM and other deep architectures like ResNet and ViT.

While these results highlight the promise of deep learning in capturing the subtle morphological features of medicinal plant leaves, we acknowledge several challenges that temper the overall significance of the study. First, the scalability of our approach remains an important consideration; the models were developed and validated on a moderately sized dataset, and their performance on larger, more diverse datasets warrants further investigation. Second,

the computational efficiency of deeper architectures such as ResNet and ViT suggests that high-performance computing resources may be required for practical, real-world deployment, which could limit accessibility for some practitioners. Lastly, while our system demonstrates high accuracy in a controlled experimental setting, the practical implementation of an automated plant identification system will need to address issues related to real-time processing, system integration, and user interface design.

References

- [1] Ramesh, S., and D. Vydeki. "Recognition and Classification of Paddy Leaf Diseases Using Optimized Deep Neural Network with Jaya Algorithm." *Information Processing in Agriculture*, vol. 7, no. 2, 2020, pp. 249–60.
- [2] Thongkhao, K., C. Tungphatthong, T. Phadungcharoen, and S. Sukrong. "The Use of Plant DNA Barcoding Coupled with HRM Analysis to Differentiate Edible Vegetables from Poisonous Plants for Food Safety." *Food Control*, vol. 109, 2020, p. 106896.
- [3] Otter, J., S. Mayer, and C.A. Tomaszewski. "Swipe Right: A Comparison of Accuracy of Plant Identification Apps for Toxic Plants." *Journal of Medical Toxicology*, vol. 17, 2021, pp. 42–47.
- [4] Alobeidli, K., M. Shatnawi, and B. Almansoori. "Image Classification for Toxic and Non-Toxic Plants in UAE." *Advances in Science and Engineering Technology International Conferences (ASET)*, 20 Feb. 2023, pp. 1–5. IEEE.
- [5] Miao, L.I., Xi-Wen Li, Bao-Sheng Li, Lu Lu, and Yue-Ying Re. "Species Identification of Poisonous Medicinal Plant Using DNA Barcoding." *Chinese Journal of Natural Medicines*, vol. 17, no. 8, 2019, pp. 585–90.
- [6] Huixian, J. "The Analysis of Plants Image Recognition Based on Deep Learning and Artificial Neural Network." *IEEE Access*, vol. 8, 2020, pp. 68828–41.
- [7] Chaudhury, A., and J.L. Barron. "Plant Species Identification from Occluded Leaf Images." *IEEE/ACM Transactions on Computational Biology and Bioinformatics*, vol. 17, no. 3, 2018, pp. 1042–55.
- [8] Zhao, Y., Z. Chen, X. Gao, W. Song, Q. Xiong, J. Hu, and Z. Zhang. "Plant Disease Detection Using Generated Leaves Based on DoubleGAN." *IEEE/ACM Transactions on Computational Biology and Bioinformatics*, vol. 19, no. 3, 2021, pp. 1817–26.
- [9] Tian, L., H. Zhang, B. Liu, J. Zhang, N. Duan, A. Yuan, and Y. Huo. "VMF-SSD: A Novel V-Space Based Multi-Scale Feature Fusion SSD for Apple Leaf Disease Detection." *IEEE/ACM Transactions on Computational Biology and Bioinformatics*, 14 Dec. 2022.
- [10] Hosny, K.M., W.M. El-Hady, F.M. Samy, E. Vrochidou, and G.A. Papakostas. "Multi-Class Classification of Plant Leaf Diseases Using Feature Fusion of Deep Convolutional Neural Network and Local Binary Pattern." *IEEE Access*, 15 Jun. 2023.
- [11] Azadnia, R., and K. Kheiralipour. "Recognition of Leaves of Different Medicinal Plant Species Using a Robust Image Processing Algorithm and Artificial Neural Networks Classifier." *Journal of Applied Research on Medicinal and Aromatic Plants*, vol. 25, 2021, p. 100327.
- [12] Wang, B., X. Yao, Y. Jiang, C. Sun, and M. Shabaz. "Design of a Real-Time Monitoring System for Smoke and Dust in Thermal Power Plants Based on Improved Genetic Algorithm." *Journal of Healthcare Engineering*, 1 Jul. 2021.
- [13] Kumar, S., A. Jain, A.P. Shukla, S. Singh, R. Raja, S. Rani, G. Harshitha, M.A. AlZain, and M. Masud. "A Comparative Analysis of Machine Learning Algorithms for Detection of Organic and Nonorganic Cotton Diseases." *Mathematical Problems in Engineering*, 16 Jun. 2021, pp. 1–8.
- [14] Poblete, T., C. Camino, P.S. Beck, A. Hornero, T. Kattenborn, M. Saponari, D. Boscia, J.A. Navas-Cortes, and P.J. Zarco-Tejada. "Detection of Xylella Fastidiosa Infection Symptoms with Airborne Multispectral and Thermal Imagery: Assessing Bandset Reduction Performance from Hyperspectral Analysis." *ISPRS Journal of Photogrammetry and Remote Sensing*, vol. 162, 2020, pp. 27–40.
- [15] Pushpanathan, K., M. Hanafi, S. Mashohor, and W.F. Fazlil Ilahi. "Machine Learning in Medicinal Plants Recognition: A Review." *Artificial Intelligence Review*, vol. 54, no. 1, 2021, pp. 305–27.
- [16] Gupta, D.P., S.H. Park, H.J. Yang, K. Suk, and G.J. Song. "Neuroprotective and Anti—Neuroinflammatory Effects of a Poisonous Plant *Croton tiglium* Linn. Extract." *Toxins*, vol. 12, no. 4, 2020, p. 261.
- [17] Prokop, P., and J. Fancovič. "The Perception of Toxic and Non-Toxic Plants by Children and Adolescents With Regard to Gender: Implications for Teaching Botany." *Journal of Biological Education*, vol. 53, no. 4, 2019, pp. 463–73.
- [18] Ketwongsa, W., S. Boonlue, and U. Kokaew. "A New Deep Learning Model for the Classification of Poisonous and Edible Mushrooms Based on Improved Alexnet Convolutional Neural Network." *Applied Sciences*, vol. 12, no. 7, 2022, p. 3409.
- [19] Hamonangan, R., M.B. Saputro, and C.B. Atmaja. "Accuracy of Classification Poisonous or Edible of Mushroom Using Naïve Bayes and K-Nearest Neighbors." *Journal of Soft Computing Exploration*, vol. 2, no. 1, 2021, p. 5360.
- [20] Cook, D., S.T. Lee, D.R. Gardner, R.J. Molyneux, R.L. Johnson, and C.M. Taylor. "Use of Herbarium Voucher Specimens to Investigate Phytochemical Composition in Poisonous Plant Research." *Journal of Agricultural and Food Chemistry*, vol. 69, no. 14, 2021, pp. 4037–47.
- [21] Panter, K.E., K.D. Welch, and D.R. Gardner. "Poisonous Plants: Biomarkers for Diagnosis." *Biomarkers in Toxicology*, 1 Jan. 2019, pp. 627–652. Academic Press.
- [22] Noor, T.H., A. Noor, and M. Elmezain. "Poisonous Plants Species Prediction Using a Convolutional Neural Network and Support Vector Machine Hybrid Model." *Electronics*, vol. 11, no. 22, 2022, p. 3690.

Authors' Profiles



Dr. Aruna S. K. is an Associate Professor in the Department of AI and Data Science Engineering, CHRIST University, Bangalore, India. She earned her PhD in Information and Communication Engineering from Anna University, Chennai. Her research focuses on Medical Image Processing, Machine Learning, and Deep Learning. She has published 25 papers in reputed international journals and has delivered several guest lectures on topics related to Computer Science and Engineering. In addition, she has completed consultancy work for Capgemini and the ICT Academy. She serves as the University Coordinator for the Honeywell Center of Excellence for Women Empowerment by ICT Academy and as the Campus Coordinator for Sub-Institutional Innovation Council (Sub-IIC) at the CHRIST–Kengeri Campus. The Ministry of MSME, Government of India, has awarded her an incubation grant of ₹21 lakhs for her startup concept in the healthcare sector.



Mr. Praveen P. is a General Manager at Sun Publications specializing in analytics and growth-driven strategies. He has a strong aptitude for using data insights to identify and capitalize on opportunities for expansion and efficiency. Mr. Praveen is dedicated to driving sustainable growth through strategic planning, data-informed decision-making, and a collaborative leadership approach. His professional focus is on optimizing performance and ensuring that Sun Publications maintains a leading position in the industry.



Mr. Gowtham K. is a Product Solution Engineer at Juspay, specializing in technical solutions. He holds a B.Tech in IT and has expertise in SQL, problem-solving, and data analysis. He has also conducted research on AI-driven plant classification and is passionate about leveraging data-driven insights for innovation in payment solutions.



Mr. Mohammed Khashif S. is a Software Engineer at Impelsys, specializing in DevOps and cloud technologies. He holds a B.Tech in IT (with Honors in Data Analytics) and has expertise in Linux, AWS, Git, and CI/CD. He has conducted research on AI-driven plant classification and has cloud experience with AWS and Terraform. He is also GitHub Foundation certified.



Keerthana Jaganathan is a researcher in the School of Computer Science at Northumbria University, Newcastle upon Tyne, UK. Her academic focus lies in the intersection of computer science and information technology, where she explores innovative solutions to real-world challenges. Keerthana's work emphasizes areas such as artificial intelligence, machine learning, and data analysis. She is committed to advancing technological research and contributing to the development of cutting-edge applications in her field. Keerthana is passionate about promoting collaborative research and education in computer science, aiming to foster both academic and industry advancements.



Dr. K. Karthick received his Ph.D. under the Faculty of Electrical Engineering from Anna University, Chennai, completing it in 2018. With over 19 years of academic experience, he currently serves as a Professor in the Department of Electrical and Electronics Engineering at GMR Institute of Technology, Rajam, Andhra Pradesh. His academic journey began with a B.E. in Electrical and Electronics Engineering from Sona College of Technology, followed by an M.E. in Power Electronics and Drives from St. Joseph's College of Engineering, Chennai. Throughout his career, He has held teaching positions at esteemed institutions such as Mahendra Engineering College, RMD Engineering College, and Panimalar Engineering College. His research interests include machine learning, power electronics, renewable energy systems, and optical character recognition. He has published extensively in renowned journals, particularly on the application of machine learning in environmental and energy-related studies. Furthermore, he is actively involved in professional organizations such as the Indian Society for Technical Education and IEEE, where he holds senior membership.

How to cite this paper: Aruna S. K., Praveen P., Gowtham K., Mohammed Khashif S., Keerthana Jaganathan, K. Karthick, "Classification of Medicinal Plant Leaves using Deep Learning Algorithms", International Journal of Intelligent Systems and Applications(IJISA), Vol.18, No.1, pp.132-149, 2026. DOI:10.5815/ijisa.2026.01.10