

Coffee Leaf Plant Disease Identification through Image Processing and Machine-Learning Techniques in Ethiopia.

Bayile Getu Taye (✉ beyagetu@gmail.com)

Indian Institute of Technology Ropar

Neeraj Goel

Indian Institute of Technology Ropar

Research Article

Keywords:

Posted Date: January 5th, 2024

DOI: <https://doi.org/10.21203/rs.3.rs-3720115/v1>

License: © ⓘ This work is licensed under a Creative Commons Attribution 4.0 International License.

[Read Full License](#)

Additional Declarations: No competing interests reported.

Coffee Leaf Plant Disease Identification through Image Processing and Machine-Learning Techniques in Ethiopia.

Authors: Bayile Getu Taye * and Neeraj Goel

Corresponding author email: beyagetu@gmail.com

Computer Science and Engineering with Specialization of Artificial Intelligence

Indian Institute of Technology Ropar

Ropar, Punjab, India

Abstract—Coffee plants are woody evergreens that can reach a height of up to ten meter's in the wild with the fruit of coffee beans, which are the seeds that produce the majority of the world's coffee. This study focuses on a variety of Arabica Coffee leaf diseases that occurs frequently in Ethiopia. The diseases which occurred primarily are Miner coffee leaf disease (MCLD), Rust coffee leaf disease (RCLD), Cercospora coffee leaf disease (CCLD), and Phoma coffee leaf disease (PCLD). The study focuses on the identification of those four a variety of diseases through image processing as well as machine learning mechanisms through Transfer learning and convolution neural network (CNN) architectures through the feed-forward model, resnet50 model, inceptionV3 model, and deep learning model through tensor flow, which are the most popular models to detect various plant diseases with high Accuracy. All the pictures used for this study were captured from the whole state of Ethiopia in every place where coffee plant exists. The total number of data sets currently used is 58,546 with 80 percent being put to use for training, while the rest 20 percent were employed for testing, with a 99.9%, 98.5%, 99%, and 99% respectively with a total success rate in classification accuracy and 99.8%, 98%, 99% and 99.7% respectively with total success rate in confusion matrix accuracy. So there is excellent performance in inferences.

I. INTRODUCTION

Coffee is grown in most parts of the world including Africa, among African countries Ethiopia produces a high amount of coffee. As a result, the agricultural industry plays a vital role in both countries' economic and social life of the people. Almost 85 percent of the population in the country is based on agricultural inputs and products. Approximately half of 85 percent of this industry with a contribution of from coffee plantation [1] [5]. As I observed for a long period, Ethiopian coffee product is the most expensive in the market among other countries' coffee product due to their being purely organic and having a uniquely delicious flavor. So every country purchased it at any time anywhere they got. Because of this reason, Ethiopia is known for the origin of coffee and it has a very high quantity and quality of coffee products. Coffee farming in the country contributed nearly a quarter of the governmental year of revenue. More than half of the population's economical life

Is based on coffee plants, either by farming or by shopping and trading in every state of the country as well as across foreign countries [2]. As we know there are different varieties of coffee in most parts of the world as well as in Africa, and Ethiopia. But Ethiopian coffee is Coffee-Arabica type which is something especial that grows within the whole states within the following proportion. Around 63 percent of the coffees are produced from Oromia state and 34 percent of the coffee in southern nation nationalities of peoples of Ethiopia with lesser amounts in Amhara state and the state of Gambelia as well as from Drie-Dawa administrative city [2]. To sum up, the majority of Ethiopian coffee-Arabica plants are grown at an elevation range of one thousand and two thousand meters. The Coffee-Arabica is an indigenous African genus (species) with several classifications found in the whole part of Africa without south and north parts of the continent [3]. Due to coffee plant infection problems as well as worldwide increasing of temperature influences, in recent years, two varieties of coffee plants are available for commercial purposes as well as economically grown in every part of the world, those are Canephora Coffee (Robusta) which grows in wetlands parts of the world and Arabica Coffee (Arabica) that are manufactured in uplands of Africa. These interesting classes of coffee_Arabica category were invented in Ethiopia, particularly in the area of Kaffa, Southern Nations, and Nationalities peoples of Ethiopia. Yemeni merchants spread Arabica coffee throughout the whole continent at the time of the fiftieth century. In the southwest and southeast of Ethiopia, less a number in tropical rain forests that yield coffee plants in a wide variety of shadow bushes (trees) [3]. The type of disease which occurs in coffee_Arabica is a disease like other Coffee Plants diseases that damages the leaf, stem, and root of coffee plants. As it showed from year to year Coffee plant infections have become a major issue in recent years which occurs around the stem, root, and coffee beans and leaves. But most of the time it occurs on leaves and its productivity decrees in quantity as well as quality, [1] [5].

II. RELATED WORKS

In research [12], the author focus on trustable predictions and metrics of a confusion matrix for rankings of web services by applying fuzzy rules to make better binary classification improved the effect by arranging the several circumstances of consumers' reaction. In research [4] [5], the authors used segmentation technique, filtering methods, and feature extraction based on inputting the infected images. Means taking as input coffee disease and Pre-processing the image. They used for minimizing noise by applying filtering methods and Classification to classify dataset with respective class. In research [6], the researchers used a sequential model as network architecture to classify datasets with the technique of Mini-Batch Gradient Descent (MBGD) which is the optimization algorithm to update parameters. Means to add or minus with a specified function and data argumentation for improvement of generalizability of the model. They also used transfer learning as a methodology for the uses of a pre-trained neural network to gain knowledge at the time of training one dataset to be used on other datasets and Categorical cross-entropy for losing functions to evaluate the performance in the model. In research [7], the authors were focused on creating a knowledge-based system by extracting knowledge of experts to make rules using decision tree methods and image processing. Images are given as input and preprocessing the image to minimize low frequency, background noise, reflection, and masking portion of the images to get good accuracy. They also used median filtering to eliminate noise and image segmentation for understanding of disease identification as a good technique. In research [4] Even if there are many techniques, they only use K-means segmentation techniques due to no single segmentation technique being appropriate to all image processing applications. They also applied feature mining (removal or extract) technique for decreasing the size of a data set by evaluating attributes to identify coffee plant disease. They also used texture features like GLCM (energy, entropy, contrast, homogeneity, correlation, shade, prominence) and Color features. In research [8], the authors used back propagation algorithm with the descending gradient method to reduce the error function by adjusting the network weights, for their entire work they evaluating different approaches based on deep learning for the problem of segmentation, classification, and quantification of biotic stress of coffee leaves. In research [9], the researchers used computational methods as like GLCM such as energy, correlation, contrast, and homogeneity as well as local binary pattern and deep learning to detect Cercospora and Rust coffee plant disease. In general in the past related works, most of the researcher works are related and the type of accuracy they got also good enough which is between 75- 97 percent, their accuracy is different according to the type of model, performance measure and dataset they have used. But they are limited with specific method as well as limited number of coffee leaf plant disease and limited number of datasets.

III. PROBLEM STATEMENT

- Predicting Coffee_Arabica leaf disease such as Miner coffee leaf disease (MCLD), Rust coffee leaf disease (RCLD), Cercospora coffee leaf disease (CCLD), and Phoma coffee leaf disease (PCLD).
- Develop and design coffee leaf plant diseases identification system through image processing as well as machine-learning mechanisms with a high accuracy.
- Measure the status of the techniques in terms of performance.

A. Evaluation Metrics

- Classification accuracy
- Confusion Matrix
 - Precision
 - Recall
 - F1-score
 - Accuracy

B. Methods

- Convolution neural network (CNN)
- Classification
- Augmentation
- Filtering
- Background removal

1) *CNN*: is a type of artificial Neural network that used to solve a high range of pattern recognition problem, including speech recognition, computer visions, coffee grading and so on [10] [11]. The primary point behind convolution neural networks is to create a model efficiently with max-pooling operation and exceptional convolution for the identification of health related issues [13]. So convolution neural networks are very fruit-full through the help of convolutional layers and other machine-learning methods.

2) *Filtering*: It learns to detect coffee disease when trained in image classifier [10].

3) *Data augmentation*: is used when there aren't sufficient pieces of training data to include a wide range of image categorization scenarios, the convolutional neural network is prone to overfitting. As a result, we employ a method so-called data augmentation to overcome the over fitting. Through numerous transformations like scaling's, rotations, and contrast_changes, data_augmentation is helped to generate fresh pieces of training of a sample from an existing piece of the training set. So by generating more data samples we can observe better performances of the model.

4) *Background removal*: it is used to reduce noise as well as used to concentrate on the infected areas, besides this, it is used to detect the front and the backside of the image.

5) *Confusion Matrix*: It is expressed in the form of a table to measure the performance as a classifier through precision, recall, and f1-score based on test data by taking true values. So to use the confusion matrix we have to know some important terms which are:-

- True positive (TP)

- False Positive(FP)
- True Negative(TN)
- False Negative(FN)

So, to elaborate on those terms let us take a Rust leaf disease-infected coffee into the consideration of binary classification.

- A coffee which was actually infected (positive) and classified as infected (positive). This is a True positive.
- A coffee which was not actually infected (negative) and classified as not infected (negative). This is a true negative.
- A coffee which was not actually infected (negative) and classified as infected (positive). This is a false positive.
- A coffee which was actually infected (positive) and classified as not infected (negative). This is a false positive.

6) *Precision*: is a performance metric that indicates the number of positive forecasts that were correct.

The formula of Precision:

$$Precision = \frac{TP}{TP + FP}$$

7) *Recall*: is a performance metric that indicates a number of the positive cases the classifier appropriately predicted. The formula of Recall:

$$Recall = \frac{TP}{TP + FN}$$

8) *F1-score*: The combination harmonic mean effect of recall and precision is called F1-score.

The formula of F1-score:

$$F1 - score = \frac{2 * Precision * Recall}{Precision + Recall} = \frac{2 * TP}{2 * TP + FP + FN}$$

The formula of Accuracy:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

C. Model Layers

1) *Keras Conv2D*: 2D Convolutions Layers which is a fundamental component of CNN architectures with a wide range of applications and produces tensors of output by winding a convolutions kernel with the layer input. Even if there are several parameters, The compulsory conv2D parameter is the number of filters from which the convolution layer will learn, which is a numeral value as well as controls the number of outputs filter in the convolutions Which follows three important steps listed below.

- K kernels waiting's till applies to the images
- Every kernel is convoluted with the volume of the inputs
- The outputs of every convoluted operation create a "2D" outputs known as an "activation map"

As everyone noted in the Fig.1 below as your result spatial

Volume is reducing your amount of filter learning will be rising

Which is a very major exercise at the time of designation of CNN architecture as well as it is better to do selecting the correct amount of filter with the power of two with value and you defiantly want to tune the correct values based on the involvedness of your datasets as well as the deepness of your NN, but it should be between ranges of [32, 64, 128] and [256, 512, 1024] from initial to the deepest layer. The other

Very important point is *kernel_size*, it depends on the inputs image, if our inputs image is higher than 128×128 we should select to use a *kernel_size* more than 3 which helps us to learn bigger spatial filter as well as to decrease the sizes of the volume. Unless use less than or equal to 3.

2) *Max Pooling*: is a feature of map merging operations that determines the highest values for patches. The features maps' dimensions are reduced by using max-pooling layers. As a result, the number of parameters to acquire (learning) and the number of images processing in the networks are both reduced. The feature contained in a section of the features map produced by convolutional layers is summed up by the max-pooling layer. See the colored table 1 below.

2	2	7	3
3	4	5	3
4	9	6	3
10	9	11	8

4	7
10	11

TABLE I: Max pooling

3) *Rectified Linear Unit*: we can't see this layer separated

From CNN due to necessary parts of its convolutional operations. The goal of using the relu's activation function is to make our pictures more non-linear. The reason behind this is that pictures are non-linear by nature.

4) *Dense layer*: is the simplest layer of a neuron in which Every neuron received inputs from all the other neurons of back layers which are used in convolutions layers outputs for classifying images.

5) *Flatten*: is the process of turning data into a one-dimensional array for use in the next layers and used to construct features of vectors which is a single lengthy output of convolution layer will be flatted, As a result, it will be connected fully with the last classifications of the model.

6) *Soft max*: is served for multi_classification problem which is beneficial since it translates the outputs of our neural network of the most end layers into what is effectively a probability distribution. So as I had five classes I used soft max at the end layer as a classifier.

$$\sigma(z_i) = \frac{e^{z_i}}{\sum_{j=1}^K e^{z_j}} \quad \text{for } i = 1, 2, \dots, K$$

Soft max for multi-class single label classification.

Sigmoid for single-label/multi-label binary classification.

$$\sigma(z) = \frac{1}{1 + e^{-z}}$$

7) *Dropout*: is not exactly a layer rather it is a method for preventing over fitting in a model.

8) *Batch-normalization*: is a method that helps to normalize contribution for the layers to each mini_batch during training in deep NN.

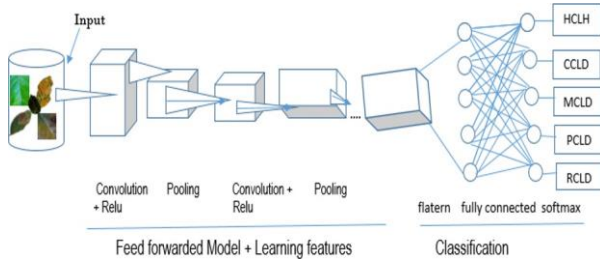


Fig. 1: Over All Architecture of feed forwarded Network.

IV. DATASET

Publically no datasets were available on this study from past researchers. Some of the datasets were available but they are less in number, has noise problem and background is not removed. There were private data sets that I had received that were taken through the expert of agricultural people using digital cameras with different sizes and needed further enhancements such as annotations with proper label names and file extensions, cropping, and other necessary preprocessing things. Those dataset were only from the leaf of Coffee Arabica type and has the overall of five classes with one class healthy Coffee leaf(HCL) and the other four classes are infected with the various coffee disease with a total number of 58,546 coffee leaf images which split into train and test in the ratio of 80 to 20 percent. Purposely I had used two types of datasets to now the effect of background noise and to check the proper inference of the model.

- Dataset with proper enhancement with a total number of 58546 images.
- Dataset without proper enhancement with a total number of 1008 images.

Here is below a sample dataset visualization for both type of datasets that we can refer in Fig 2 and 3 respectively.

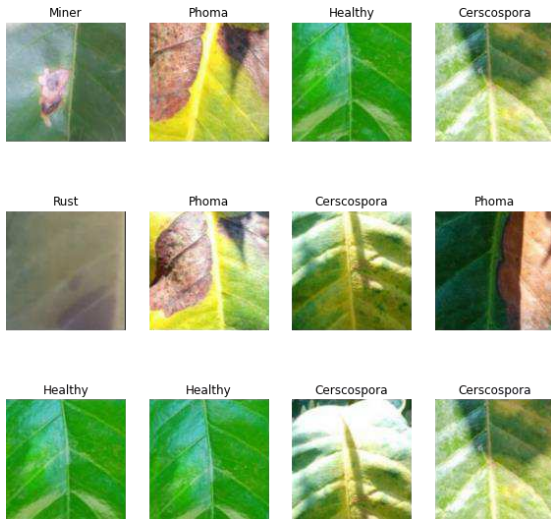


Fig. 2: sample dataset visualization proper enhancement.

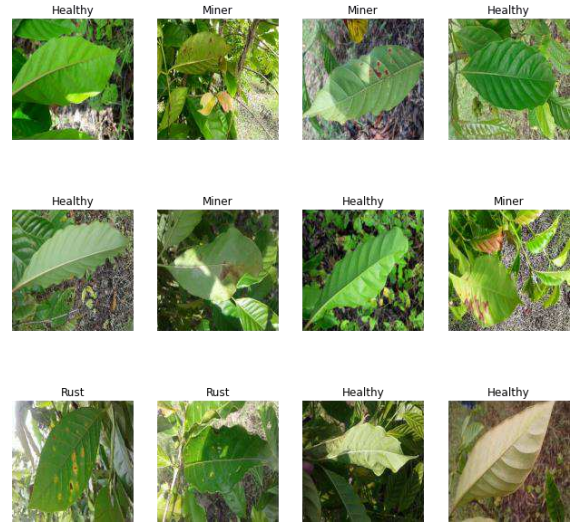


Fig. 3: sample dataset visualization without enhancement.

A. Class wise distribution of dataset images

The number of times each class appears in each image in a dataset is quantified by the Distribution of Class Frequency per Image. This statistic shows the frequency and distribution of various classes across the dataset, illuminating the class heterogeneity among various photos. For results refer Fig. 4

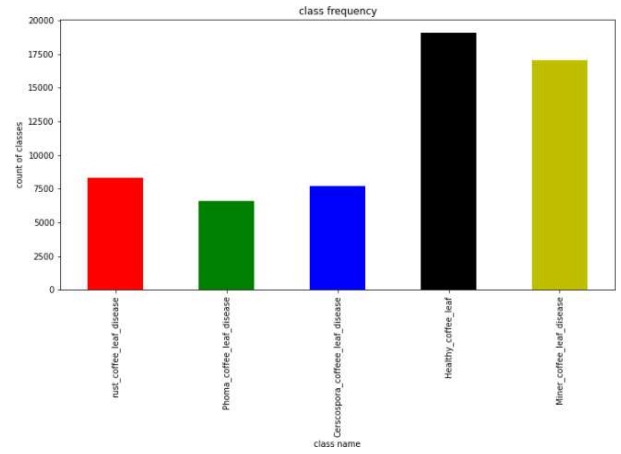


Fig. 4: Class wise distribution of dataset images.

V. EXPERIMENT

Experiment is done through four models namely the feed forwarded model using CNN, Resnet50 model, InceptionV3 model, and Deep learning model through the tensor flow to predict and develop coffee leaf plant disease through image processing as well as machine-learning mechanisms. So in this experiment, I used both the performance measure metrics such as classification accuracy and confusion matrix accuracy.

1) *Feed forwarded Model:* In this model, each input was entered into the layers. Then multiplied by weights and every value should be added together to obtain a summation of the weighted input values. For results refer Fig. 5, 6 and 7

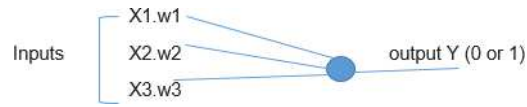


Fig. 5: How Feed Forward NN works.

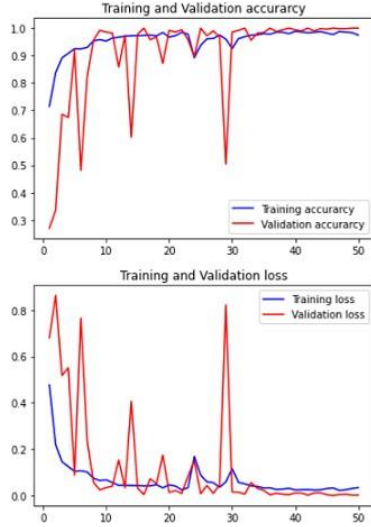


Fig. 6: Training and Validation Accuracy and loss respectively.

Actual Classes	Predicted Classes				
	CCLD	RCLD	PCLD	HCL	MCLD
CCLD	108	0	0	0	0
RCLD	0	111	0	0	0
PCLD	0	0	93	0	0
HCL	0	0	0	88	0
MCLD	0	0	0	0	100

TABLE II: Confusion Matrix for feed forward model

	precision	recall	F1-score	Support
0	1.00	1.00	1.00	108
1	1.00	1.00	1.00	111
2	1.00	1.00	1.00	93
3	1.00	1.00	1.00	88
4	1.00	1.00	1.00	100
accuracy			1.00	500
macro avg	1.00	1.00	1.00	500
weighted avg	1.00	1.00	1.00	500

TABLE III: Classification Report for feed forward model

2) *Resnet50 model*: It is CNN that is fifty layers deep and influential backbone model used for classifications tasks. It uses to miss connection to add the result from a former layer to a future layer. This supports it to diminish the vanishing gradient problem. For results refer Fig. 8 and 9

3) *InceptionV3 model*: This model is used for image classification through deep learning based on CNN. For results refer Fig. 10 and 11

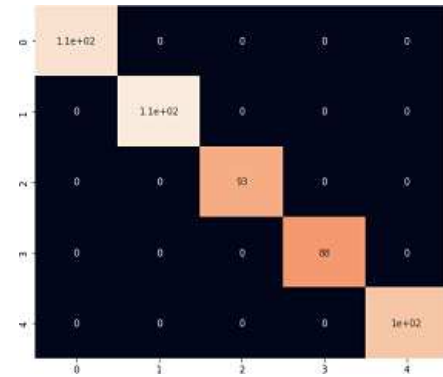


Fig. 7: visualization with seaborn for feed forward model.

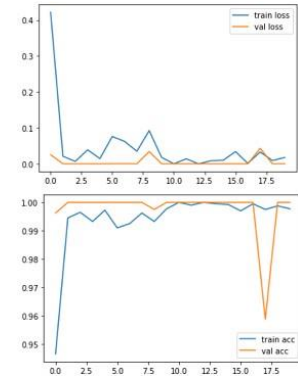


Fig. 8: Training and Validation loss and Accuracy respectively.

Actual Classes	Predicted Classes				
	CCLD	RCLD	PCLD	HCL	MCLD
CCLD	36	0	0	0	1
RCLD	0	38	0	0	0
PCLD	0	0	44	0	0
HCL	0	0	2	46	0
MCLD	1	0	0	0	32

TABLE IV: Confusion Matrix for Resnet 50 model

	precision	recall	F1-score	Support
0	0.97	0.97	0.97	37
1	1.00	1.00	1.00	38
2	0.96	1.00	0.98	44
3	1.00	0.96	0.98	48
4	0.97	0.97	0.97	33
accuracy			0.98	200
macro avg	0.98	0.98	0.98	200
weighted avg	0.98	0.98	0.98	200

TABLE V: Classification Report for Resnet50 model

Actual Classes	Predicted Classes				
	CCLD	RCLD	PCLD	HCL	MCLD
CCLD	47	0	1	1	0
RCLD	0	66	0	0	0
PCLD	0	0	64	0	0
HCL	2	0	0	63	0
MCLD	0	0	0	0	56

TABLE VI: Confusion Matrix for InceptionV3 model

4) *Deep learning model*: Tasks such as text classification, sound classification or image classification are executed by

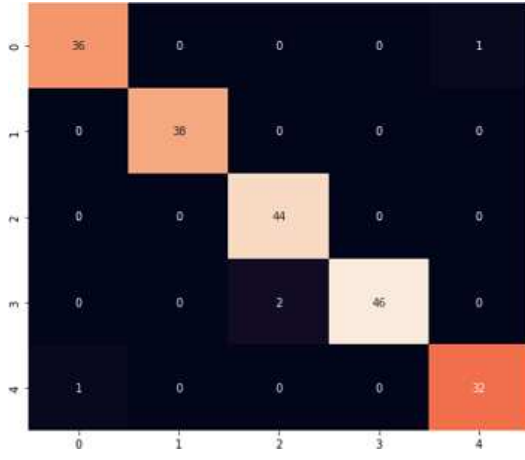


Fig. 9: visualization with seaborn for Resnet50 model.

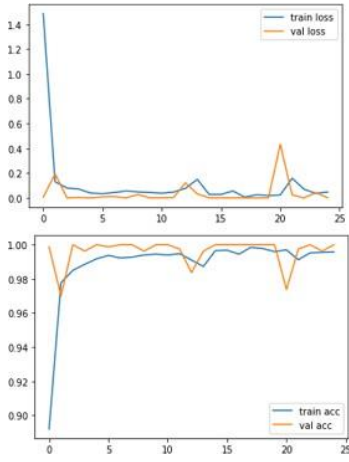


Fig. 10: Training and Validation loss and Accuracy respectively.

	precision	recall	F1-score	Support
0	0.96	0.96	0.96	49
1	1.00	1.00	1.00	66
2	0.98	1.00	0.99	64
3	0.98	0.97	0.98	65
4	1.00	1.00	1.00	56
accuracy			0.99	300
macro avg	0.99	0.99	0.99	300
weighted avg	0.99	0.99	0.99	300

TABLE VII: Classification Report for InceptionV3 model

Computer model through learning is called deep learning. So in this model we can reach state-of-the-art accuracy. Even we may get more than human beings performance. For results refer Fig. 12 and 13

Actual Classes	Predicted Classes				
	CCLD	RCLD	PCLD	HCL	MCLD
CCLD	67	0	1	0	0
RCLD	0	77	0	0	0
PCLD	0	0	90	0	0
HCL	0	0	0	82	0
MCLD	0	0	0	0	83

TABLE VIII: Confusion Matrix for Deep Learning model

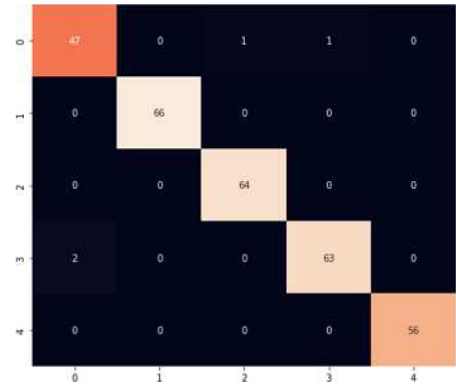


Fig. 11: visualization with seaborn for InceptionV3 model.

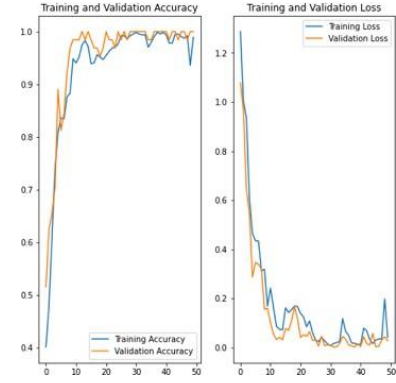


Fig. 12: Training and Validation loss and Accuracy respectively.

	precision	recall	F1-score	Support
0	1.00	0.99	0.99	68
1	1.00	1.00	1.00	67
2	0.99	1.00	0.99	90
3	1.00	1.00	1.00	82
4	1.00	1.00	1.00	83
accuracy			1.00	400
macro avg	1.00	1.00	1.00	400
weighted avg	1.00	1.00	1.00	400

TABLE IX: Classification Report for Deep learning model

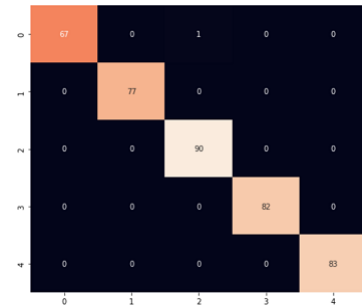


Fig. 13: visualization with seaborn for Deep learning model.

VI. INFERENCE MODEL

Inference is the process of Operating a machine learning model or bringing a machine learning model into production which means predictions of the model. So in this study

Inference is done with in two different data points namely dataset with proper enhancement with 58,546 images and dataset without proper enhancement with 1008 images to check the performance of models in various data points. For results refer Fig. 14, 15, 17 and 18

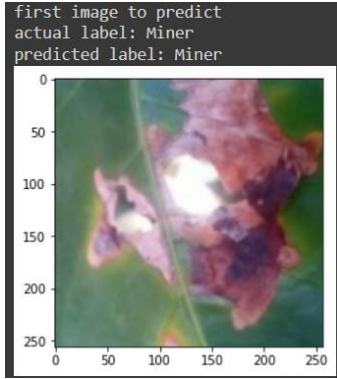


Fig. 14: first image to predict.

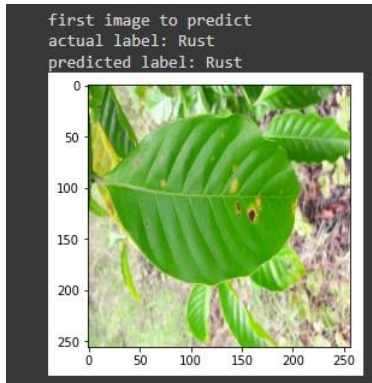


Fig. 15: first image to predict with less enhanced data points.

A. Probability Distribution

It is a function that shows the probabilities of each likely event for a given random variable. To put it another way, it describes the range of possible event frequencies. For results refer Fig. 16

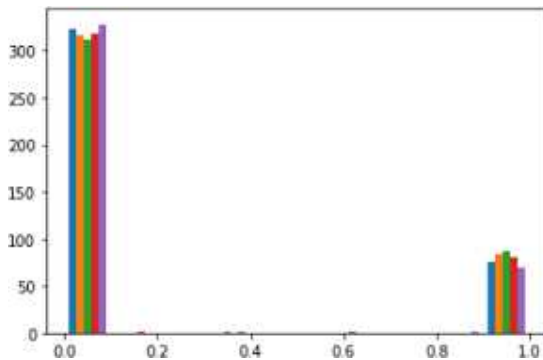


Fig. 16: Probability Distribution.

VII. CONCLUSIONS AND COMPARISON OF MODELS

To conclude the performance metrics of confusion matrix accuracy are much better than the performance metrics of classification accuracy and among the four models used in this study feed forwarded model achieves 99.8% accuracy. Another core idea in this study is, that if we have imbalanced data points then the confusion matrix is the best one. For more see the table below named a Comparison of models.

VIII. DECLARATIONS

A. Ethics approval and consent to participate

The Indian Institute of Technology Roppar's ethical clearance review committee gave their approval to this work. Following their disclosure of the research's overall goal, the Ethiopian coffee research also indicated their desire to participate. All research subjects provided written consent. The goal of the study and the participants' right to decline to answer questions for the data collectors were explained to them. Confidentiality was maintained for all data collected for the study.

B. Consent for publication

We, B G T and N G, hereby declare that:

- 1) This paper's research was carried out in compliance with ethical standards and guidelines.
- 2) All persons engaged in the research, including participants, co-authors, and any other pertinent parties, have given their informed consent.
- 3) The participants have given their informed consent to participate in the study, having been informed of its nature, goal, and possible risks and benefits.
- 4) All sensitive or private data gathered for the study has been treated with the highest confidentiality and will be presented in a way that protects the participants' privacy and anonymity.
- 5) Any use of copyrighted materials, such as pictures, figures, or data, has been approved and authorized by all relevant parties, and appropriate credit has been given.
- 6) We acknowledge that getting consent and upholding ethical behavior are the researchers' responsibilities, and we accept full responsibility for any problems that may result from this research's dissemination.
- 7) We understand that this publication may be withdrawn or retracted in the event that ethical standards are broken or appropriate consent is not obtained.

C. Availability of data and materials

We hereby declare that the dataset used in our research is available for access and use by interested researchers. The dataset is currently stored locally on our computer. To request access to the dataset, please contact us at beyagetu@gmail.com. We are happy to share the dataset with individuals who are interested in utilizing it for their research purposes. Please note that the dataset is provided under the following terms and conditions: such as non-commercial use



Fig. 17: sample inferences of the model.



Fig. 18: sample inferences of the model with less enhanced data points .

models	Classification accuracy	Confusion matrix accuracy	Epochs	No of data points	Macro avg	Weighted avg	Support
Feed forward	99.9	99.8	50	58546	100	100	500
Resnet 50	98.5	98	20	58546	98	98	200
Inception V3	99	99	25	58546	99	99	300
Deep Learning	99	99.7	50	58546	100	100	400

TABLE X: Comparison of models

or proper attribution. Furthermore, we are open to collaboration opportunities with researchers who wish to explore further research avenues using the dataset.

D. Competing interests

The authors of this article declare that they have no competing interests.

E. Funding

While doing this research, we were paid a monthly stipend for about a year by the Indian Institute of Technology, Ropar. Although we are paid a stipend, there is no interference or influence in our research work. Besides this, Based on your publication policies as it states that,” Springer Nature offers APC waivers to papers whose corresponding authors are based in countries classified by the World Bank as low-income economies as of July 2023.” as a result the corresponding author is from Ethiopia which is classified as low-income economies. So we kindly request you to consider this issue.

F. Authors’ contributions

The article’s conception and design, data collection, analysis, and interpretation were significantly influenced by the authors Bayile Getu Taye and Neeraj Goel. The authors Bayile Getu Taye and Neeraj Goel also agreed to take responsibility for all aspects of the work, participated in its writing or critical revision for significant intellectual content, and approved the final version before it was published. The final manuscript was read and approved by all writers (Bayile Getu Taye and Neeraj Goel).

G. Acknowledgment

This MTP could not be accomplished with the effort of a few people; many have Contributed their best for the successful ongoing of this MTP. First and foremost I Bring the greatest gratitude to my Supernatural God for his ceaseless help throughout my Whole works. Secondly, I offer my sincere thanks to my guide Dr. Neeraj Goel (Assistance professor, Ph.D.), IIT Ropar, Punjab, India for his constrictive comments, guidance, and help in every step and Mr. Hasabu who help me by providing detests, Last but not least, my heartfelt thanks and respect to all those people who were not mentioned here but their contributions have been inspiring me for ongoing of this MTP.

REFERENCES

- [1] S. N. Ghaiwat and P. Arora, “Detection and Classification of Coffee diseases Using Image processing Techniques”, A Review, vol. 2, no. 3, (2014).
- [2] URL: http://www.apsnet.org/edcenter/intropp/topics/Pages/Plant_Disease_Diagnosis.aspx.
- [3] S. B. Dhaygude and N. P. Kumbhar, “Agricultural Coffee diseases Disease Detection Using Image Processing”, vol. 2, no. 1, (2013).
- [4] A. Salem, B. Samma and R. A. Salam, “Adaptation of K-Means Algorithm for Image Segmentation”, CiteSeerx.
- [5] Abraham Debasu Mengistu, Dagnachew Melesew Alemayehu, and Seffi Gebeyehu Mengistu. Ethiopian coffee plant disease recognition based on imaging and machine learning techniques. International Journal of Database Theory and Application, 9(4):79–88, 2016.
- [6] Manoj Kumar, Pranav Gupta, Puneet Madhav, et al. Disease detection in coffee plants using a convolutional neural network. In 2020 5th International Conference on Communication and Electronics Systems (ICCES), pages 755–760. IEEE, 2020.
- [7] Abraham Debasu Mengistu, Seffi Gebeyehu Mengistu, and DM Alemayehu. An automatic coffee plant disease identification using hybrid approaches of image processing and decision tree. Indonesian J.Electr. Eng. Comput. Sci, 9(3):806–811, 2018.
- [8] Abraham Debasu Mengistu, Seffi Gebeyehu Mengistu, and DM Alemayehu. An automatic coffee plant disease identification using hybrid approaches of image processing and decision tree. Indonesian J.Electr. Eng. Comput. Sci, 9(3):806–811, 2018.
- [9] Lucas Ximenes Boa Sorte, Carolina Toledo Ferraz, Francisco Fambrini, Roseli dos Reis Goulart, and Jose’ Hiroki Saito. Coffee leaf disease recognition is based on deep learning and texture attributes. Procedia Computer Science, 159:135–144, 2019.
- [10] Goodfellow, Ian, and Bengio, Yoshua and Courville, Aaron. Machine learning basics. MIT Press Cambridge, MA, USA. Deep learning, 7(1):98–164, 2016.
- [11] Walleign, Serawork and Polceanu, Mihai and Jemal, Towfik and Buche, C édrick. Coffee Grading with Convolutional Neural Networks using Small Datasets with High Variance. V ’a clav Skala-UNION Agency, 2019.
- [12] Hasnain, Muhammad and Pasha, Muhammad Fermi and Ghani, Imran and Imran, Muhammad and Alzahrani, Mohammed Y and Budiarto, Rahmat, Evaluating trust prediction and confusion matrix measures for web services ranking, IEEE Access, 8:90847–90861, 2020.
- [13] Huang, Shih-Chia and Le, Trung-Hieu, Principles and Labs for Deep Learning, Academic Press, 2021