

Automated Quality Assessment using Computer Vision and Data-Driven Techniques

Anshul Bansal¹, Dhruv Gangwani², Akshita Singh³ and Nepali Singla⁴

¹⁻⁴Department of Computer Science, ABES Engineering College, Ghaziabad, U.P., India

Email: anshul.21b1541124@abes.ac.in, dhruv.21b0121191@abes.ac.in, akshita.21b0121070@abes.ac.in, nepali.singla@abes.ac.in

Abstract— Traditional crop quality grading, which is done by visual observation by humans, is by nature not perfect. Human fatigue, subjectivity, and differing degrees of experience lead to inconsistent grading, having negative impacts on farmers' revenues and consumer confidence. The technology of computer vision, machine learning, and artificial intelligence has radically transformed this field, offering automated systems for objective and precise crop monitoring. These systems scan vast data sets, detecting subtle defects in quality and disease at early stages that elude human inspectors. Free from the limitations of human judgment, these technologies provide continuous, objective decisions, providing consistent and precise grading. In real time, data on plant health, growth patterns, and yield potential enable farmers to make informed irrigation, fertilization, and pest control decisions.

Along with that, automated systems also offer early stress factor detection to allow for timely interventions and losses to be prevented. Such proactive intervention stimulates sustainable agriculture by utilizing the available resources optimally and preventing unnecessary uses of chemicals. In the end, such state-of-the-art technologies lead to a more resilient, efficient, and profitable agriculture system for farmers and consumers through high-quality consistent crops.

Keywords— Crop quality assessment, Machine learning, Sensing technologies, Data on crop health, size, and maturity, Automation, Efficient resource allocation, Increased efficiency, Reduced post-harvest losses, Constant product quality, Automated systems, Precision agriculture techniques, Prioritizing interventions, Data-driven, crop management.

I. INTRODUCTION

Traditional farming is heavily dependent on manual Labor and human judgment. This creates problems in getting the best crop quality and quantity. Mistakes by people, changes in weather, and pest attacks can greatly affect farm productivity [5]. Solving these problems with advanced technologies such as computer vision and artificial intelligence can be of great potential [16].

This project will design a strong and effective computer vision system that may inspect crop quality very fast. The system will monitor images taken from various sources, such as drones or field cameras, in order to identify important features and will sort and rate crops through the use of machine learning methods in accordance with established standards for quality [20].

A. Primary Goals

Better Quality Control: The system will make the quality checking process automatic, hence making grading steady and correct. The system will raise product quality and lower losses after harvesting due to the eradication of human mistakes [19].

Better Utilization of Resources: Farmers can monitor how their crops are doing and how they are growing in real-time. This will help them use resources like water and fertilizers more efficiently. This method of farming will reduce waste and increase harvest amounts[17].

Better Decision-Making: With the information from the system, the farmer is going to make smarter decisions about planting, harvesting, and what to do with the harvest. Using past data and guessing future patterns by farmers will improve their ways and reduce risks.

More Sustainability: Proper utilization of available resources with a minimum intake of harmful chemicals will be there in this system, and that will support sustainable farming. This is how it results in conserving the environment also, which makes the practice successful in the long run too[1].

B. Technical Method

The proposed system will consist of the following main components:

Image Acquisition: Images of crops will be taken with high resolution from drones, ground-based cameras, and satellite imagery.

Image Preprocessing: Images captured would be processed with a view to enhance their quality and eliminate noise and other artifacts. The methods used would include noise removal, image sharpening, colour correction, etc.

Feature Extraction: Colour, texture, and shape information along with patterns will be extracted from the prepared images. Feature vectors are needed for this kind of classification and regression[13].

Teaching Models using Machine Learning Models Advanced machine learning methods like CNNs will be used with large sets of images with labels to teach these machines to understand patterns and, consequently, sort crops by quality factors[7].

Real-time Analysis: The trained models would be deployed on edge devices or cloud platforms, which would allow real-time analysis of incoming images. This would give quick feedback to farmers and agricultural experts[18].

Expected Results: Accurate Crop Quality Attributes: The system will accurately determine the crop quality attributes that may include ripeness, disease, and pest infestation[8].

Optimized Resource Utilization: Real-time data will enable farmers to optimize the utilization of resources such as water and fertilizers.

Better Decision-Making: Information from data will help farmers make smart choices about planting, harvesting, and what to do after harvesting[15].

Greater Sustainability: The system will increase sustainability in farming, as it is not only reducing dependence on chemicals and reducing damage to the environment, but also giving the power back to farmers to efficiently use resources and reduce waste through automation and knowledge of crop health in real-time. This model creates an entirely new agricultural economy based on data informed decision making that will improve yields, while conserving natural resources; With the predictions of the future growth in human population for the next few decades, it is clear to see that the project is fundamentally about building a more sustainable and resilient food system moving forward, in the long-term [18].

II. METHODOLOGY

The methodology for the project follows a systematic process to ensure comprehensive and accurate analysis:

- **Data Collection:** The first step is gathering a comprehensive set of data on each student. This data consists of student variables that could impact student success (e.g., attendance data, time spent on study, parents' level of education, socio-economic status, engagement with resources). Data is sourced from many sources such as surveys, student information systems, and external sourced educational data.
- **Data Preprocessing:** Raw data tends to be messy and inconsistent. During data preprocessing, we will address the missing values in the data and implement missing data imputation methods such as mean or median imputation, and removing outliers, deal with categorical variables in numerical format using one-hot encoding, etc. This process is aimed to produce the data in a "clean" state for regression analysis and it is also compatible with machine learning algorithms, and potentially improves accuracy for prediction.

- **Model Development:** Next, regression models will be constructed to understand how the predictors are associated with student outcomes. Initially, we will construct simple linear regression models to examine individual predictors, and then multiple linear regression models containing multiple factors at a time. At this point Student factors such as study time, parental education and attendance can be quantified in terms of impact on educational outcomes.
- **Validation:** Once we develop the models, we validate their performance on some critical metrics to assess their reliability and predictive strength. The primary metrics of assessment are R-squared (percentage of variance accounted for by the model), Mean Absolute Error (MAE), and Root Mean Squared Error (RMSE). These measurements enable us to know the correctness of the predictions the model is making as well as the degree to which the models may be tuned for improvement.
- **Visualization:** After the regression models have been tested, the findings are reported in simple and visually appealing charts, graphs, and tables. This makes it easy for educators and administrators to interpret the results. Visualization tools assist in bringing out key predictors of student performance, spotting trends in the data, and demonstrating the impact of various interventions or strategies.

A. System Architecture

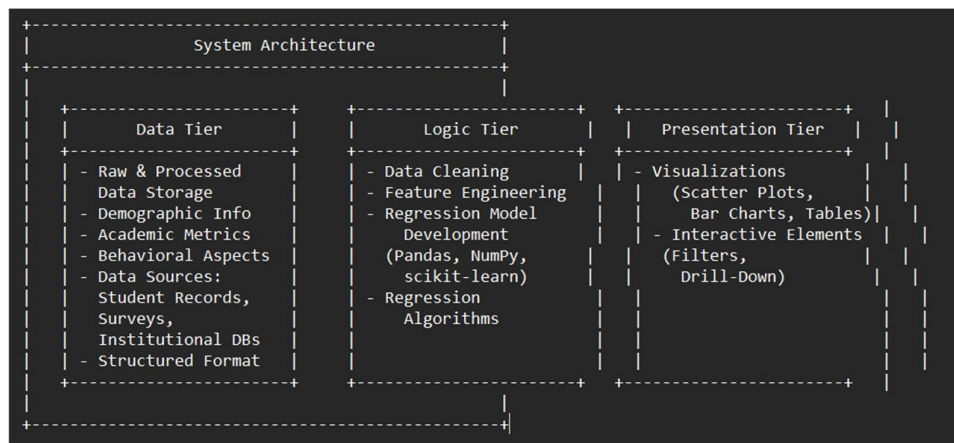


Figure 1: System Architecture

The system is designed with three-tier architecture to ensure efficient data processing, model development, and presentation of results:

- **Data Tier:** This level holds raw and processed data of students. It has a high number of variables such as demographic information, measures of academic performance (e.g., grades, test scores), and behavioral measures (e.g., study time, attendance). It also executes preprocessed data, which is cleaned and in a form ready for analysis. Raw data are collected from diverse sources such as student records, surveys, and institutional databases. The information is therefore stored in a structured manner for maintenance and easy retrieval.
- **Logic Tier:** The logic tier is the core of the system, where regression algorithms are executed to process data. Python libraries such as Pandas, NumPy, and scikit-learn are utilized to process and analyze the data. Data cleaning, feature engineering, and creation of regression models are performed in the tier. Regression methods like simple and multiple linear regression are used to analyze the relationships between grades and other factors like attendance, study hours, socio-economic status, and parental education level. Data preprocessing activities like missing value replacement and encoding categorical variables to set the dataset ready for use in machine learning fall under this level.
- **Presentation Tier:** The presentation tier is designed to communicate the results of the analysis in a way that is interpretable and understandable by teachers, administrators, and stakeholders. It produces easy-to-interpret visualizations such as scatter plots, bar graphs, and tables to display regression coefficients, model performance, and the relationship between various predictors and student outcomes. This tier also includes interactive elements, such as filters and drill-down, to allow users to drill deeper into the data and outcomes. Presentation layer ensures that key information is presented clearly, allowing educational institutions to make informed decisions to improve student performance based on data.

III. IMPLEMENTATION PLAN



Figure 2: Implementation Plan

A detailed data flow chart (Figure b) provides a visualization of the entire process, from the initial data-gathering phase to the final results presentation. The chart describes how data moves across the system, referencing important phases including data gathering, preprocessing, regression analysis, and visualization. It shows the way raw data is converted at every stage so that stakeholders can trace the process from data input to meaningful conclusions.

A. Data Collection

Data collection is the basis of every data undertaking, where raw data is being unearthed from different sources such as databases, files, web scraping, sensors, and surveys. Determining data requirements, data sources related to them, and selecting relevant modes of collection must be accomplished. Data governance and ethical standards should be applied so that data management can be rightly governed.

B. Preprocessing

Preprocessing of data is an important task that converts raw data into ready form. Cleaning is done while preprocessing, and it fills the missing values, removes duplicates, and corrects them as well. Normalization converts data in order to prepare it for proper analysis. Preprocessing is a very important procedure for correct and reliable results.

C. Feature Extraction

Feature selection is the method through which one selects the optimal features of preprocessed data in order to develop a model. One can achieve this by creating new features or modifying the current ones. It minimizes complexity, enhances the quality of the model, and uncovers concealed patterns.

D. Quality Assessment

It is crucial to keep data quality intact while analyzing. It keeps all things correct and meaningful. You check whether data is complete, consistent, accurate, and appropriate. This will help you identify and rectify errors to keep data used in analysis accurate and meaningful.

E. Post-processing and Analysis

In this phase, various methods are utilized to analyze the prepared data. These may include the application of statistical techniques in order to find patterns, machine learning to identify trends, or creating visual charts and graphs in order to reveal the data further. Once these techniques have been applied, results are analyzed to understand their meaning. This activity helps in forming conclusions or predicting future events on the basis of provided information.

F. Deployment

Here, you ready the model or analysis for actual usage. This may include integrating the model into a computer program to implement correctly, building a dashboard to display data properly, or creating detailed reports to

encapsulate significant findings and results. Your objective is to make sure information is well communicated to the stakeholders.

G. User Interaction

In this final stage, we make it possible for individuals to use the system or analysis we've created. We do this by putting in user interfaces, that is, the buttons and screens they use, and by establishing APIs, which allow other computer programs to communicate with each other. We also implement feedback systems to allow users to comment or report problems. These features make the system easy to use and encourage fruitful interaction.

IV. IMPLEMENTATION AND RESULTS

A. Tools and Technologies

Implementation of this project was based on Python and its powerful library culture, which facilitates gigantic functionalities to handle data manipulation, machine learning, and data visualization. Key technologies and tools used were:

- Pandas: For data cleaning, preprocessing, and manipulation.
- NumPy: For numerical operations and efficient data handling.
- scikit-learn: For building and evaluating regression models.
- Matplotlib and Seaborn: For visualizing relationships between variables and presenting the results in graphical formats.
- Jupyter Notebook: As the integrated development environment (IDE) to execute and document the analysis interactively.

Additionally, hardware specifications included an Intel Core i5 processor, 8GB RAM, and 256GB SSD, ensuring smooth execution of data processing tasks and model training.

B. Random forest model

Random Forest is a cross-platform machine learning algorithm of unparalleled popularity grounded in correctness, reliability, and interpretability. Random Forest is one of the group of ensemble learning algorithms that depends on many models for boosted performance.

Random Forest operates on the premise that it trains a large number of decision trees and averages the predictions of outcomes. Merging several decisions enables the algorithm to minimize overfitting and increase generalizability.

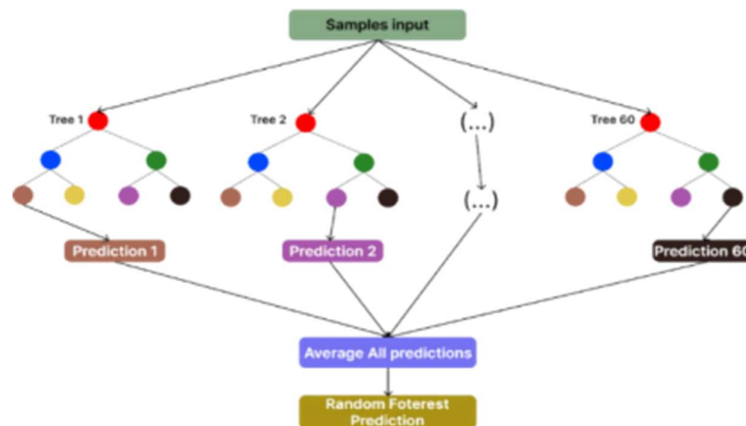


Figure 3: Working of Random Forest Model

This is the overall concept behind the Random Forest algorithm, which can be understood from the image: The above image depicts that several decision trees are learned on a random subset of the input data and then several trees are combined in order to generate the final prediction. This model of ensemble where several models are combined can exploit the power of random forest fully. With the collective opinion of diverse trees, Random Forest can effectively reduce overfitting through its improvement in accuracy as it tries to tackle complex relationships in the data.

Algorithm

Let the original dataset be defined by $D=\{x_1,x_2,\dots,x_n\}$, where n is the number of data points.

Now we generate a bootstrap sample $D_{\text{sample}}=\{x_{i1},x_{i2},\dots,x_{im}\}$ with the sample size m normally equal to n , the size of the original dataset.

The sampling process can be stated in the following way:

$$ij \sim \text{Uniform}(1,n), \text{ for } j=1,2,\dots,m$$

Where:

- ij is the index of the data point selected in the j -th draw.
- $\text{Uniform}(1,n)$ says that each index ij is randomly selected from the set $\{1,2,\dots,n\}$ with equal probability.
- m is the size of the bootstrap sample.

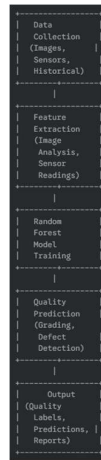


Figure 4. Random Forest Integration in Crop Quality Assessment

This image illustrates the usage of Random Forest in predicting crop quality. It displays data collection, preprocessing, feature extraction, and training with numerous decision trees. The model predicts and evaluates crop quality.

C. Resnet

ResNet changed how deep learning is done by allowing very deep neural networks to be trained effectively. It tackles the "degradation problem" with features called "residual blocks" and "skip connections." These features let the network layers learn "residual mappings," which means they can pass information through without altering it. This improves the flow of gradients and increases accuracy in tasks like recognizing images. The concepts behind ResNet are now applied broadly in many areas.

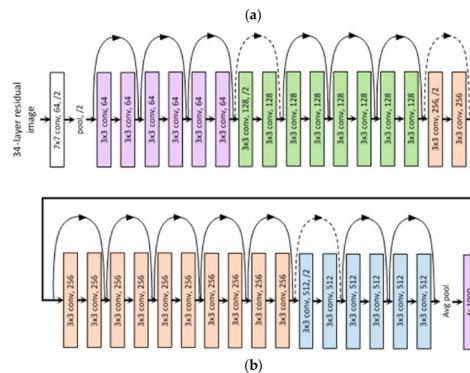


Figure 5: Working of Random Forest Model

This image illustrates two types of building blocks used in ResNet models, which are a form of neural networks. On the left, there is a basic residual block, while on the right, there is a more complex block with a 1x1 convolution to align sizes. Both blocks have layers that perform convolutions, similar to filters to look for vital information in data. They consist of a ReLU activation function to introduce non-linearity, which allows for learning, and batch normalization to enhance the learning speed and stability. One intriguing feature of these blocks is skip connections, which are shortcuts to improve learning efficiency and training deeper layers. These are the required components to enable the deep network to be capable of understanding complex patterns without encountering difficulties such as plateauing or slowing down.



Figure 6: ResNet Integration in Crop Quality Assessment

This image shows the role of ResNet in crop quality inspection. Crop images are captured, preprocessed, and assessed using the ResNet model. ResNet is fed features to assess crop health, maturity, and detect faults. The system provides accurate, automatic assessments for better farming decisions.

The output of the ResNet model is further utilized to determine crop quality. This evaluation may be in different forms, including:

- Categorization of crops into various quality grades.
- Identification of certain problems such as illnesses or nutritional deficiencies.
- Quantitative measures of quality parameters.

D. Decision Tree Construction

For every training set, a decision tree is built. At every node of the tree, it selects a random subset of features for splitting the data. Best feature and threshold are chosen based on a criterion such as information gain or Gini impurity.

This process continues recursively until some stopping criterion is satisfied (for example, maximum depth or minimum number of samples).

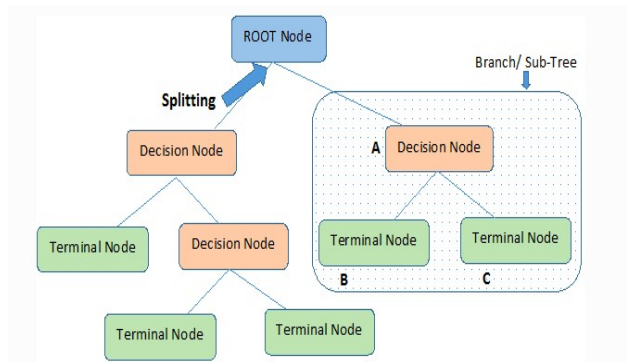


Figure 7: Working of Decision Tree Construction

Decision trees are a method used in supervised learning for tasks such as classification and regression. Picture them as actual trees, but with branches and nodes that make decisions based on certain conditions to predict results.

They are straightforward and easy to grasp, making them quite popular. However, they can occasionally become too detailed, which leads to a problem called overfitting. Despite this issue, decision trees are extensively used across various fields for making predictions and supporting decision-making processes.

E. Ensemble Learning

Final prediction is done by combining the predictions of all decision trees. For classification problems, the class which occurs most often among the predictions is chosen.

For regression tasks, the mean of the predictions is used.

F. Data Cleaning and Preprocessing

Data preparation in the correct way was responsible for maintaining healthy results. Raw data contained student demographic information, educational background, and socio-economic information. Pre-processing involved:

- Dealing with Missing Values: The missing values were imputed based on methods like mean substitution in the event of numerical data and mode substitution in the event of categorical data.
- Outlier Removal and Detection: Statistical procedures (e.g., Z-scores and IQR) were used to detect outliers, and they were either removed or transformed to prevent distorting the results.
- Feature Encoding: Parental education level and type of school, being categorical variables, were converted into numerical variables using one-hot encoding and label encoding for use in regression models.
- Normalization and Standardization: Continuous variables were normalized for uniformity and to prevent domination by the variables with larger ranges.

G. Feature Selection

A critical aspect of the analysis was identifying the most significant predictors of student academic performance. Feature selection was guided by:

- Correlation Analysis: Pearson and Spearman correlation coefficients were employed to test the strength and direction of association between independent variables (e.g., study time, attendance) and the dependent variable (performance).
- Variance Inflation Factor (VIF): Used to check multicollinearity among predictors and eliminate redundant variables.
- Domain Knowledge: Insights from prior studies and expert input informed the inclusion of variables such as parental education and socio-economic status.

The final set of predictors included attendance, study hours, parental education level, extra-curricular involvement, and socio-economic background.

H. Regression Model Development

Two primary regression techniques were implemented to model student performance:

- Simple Linear Regression: Focused on individual predictors (e.g., attendance) to assess their standalone impact on academic outcomes.
- Multiple Linear Regression: Incorporated all selected predictors to evaluate their combined effect on student performance.

The steps for model development included:

- Model Training: The dataset was split into training (80%) and testing (20%) sets to ensure unbiased evaluation.
- Model Fitting: Regression models were trained on the training set using scikit-learn's implementation of linear regression.
- Model Evaluation: The models were evaluated using metrics such as:
 - R-squared: To determine the proportion of variance in the dependent variable explained by the predictors.
 - Mean Absolute Error (MAE) and Root Mean Squared Error (RMSE): To quantify prediction errors and assess model accuracy.

I. Results and Interpretation

The regression models provided a detailed understanding of the relationships between predictors and academic performance. Key findings included:

- Attendance and Study Time: Both variables were positively correlated with performance, with study time having a slightly stronger impact.
- Parental Education: Higher levels of parental education significantly contributed to better academic outcomes, underscoring the role of a supportive home environment.
- Socio-Economic Status: Students from higher socio-economic backgrounds tended to perform better, although this variable showed moderate multicollinearity with parental education.

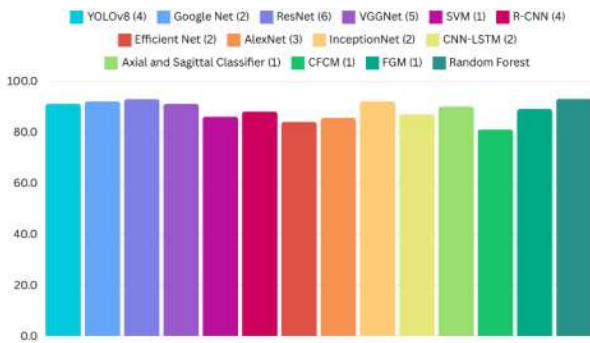


Figure 8: Comparison of Model Performance

The chart shows that models like YOLOv8, ResNet, and Random Forest have high scores (over 80.0), whereas others like SVM, CFCM, and FGM have low scores. This suggests that in the specific task and dataset used in this evaluation, YOLOv8, ResNet, and Random Forest are superior.

Y-Axis: Shows the performance measure, presumably accuracy or an equivalent value, between 0.0 and 100.0. The axis is marked with a step of 20.0.

X-Axis: Shows various machine learning and deep learning models. The models are grouped as follows:

Object Detection: YOLOv8 (4)

Convolutional Neural Networks (CNNs):

GoogleNet (2)

ResNet (6)

VGGNet (5)

InceptionNet (2)

EfficientNet (2)

AlexNet (3)

Other Models:

SVM (Support Vector Machine) (1)

R-CNN (Region-based Convolutional Neural Network) (4)

CNN-LSTM (Convolutional Neural Network - Long Short-Term Memory) (2)

Axial and Sagittal Classifier (1)

CFCM (1)

FGM (1)

Random Forest

V. IMPLEMENTATION AND RESULTS

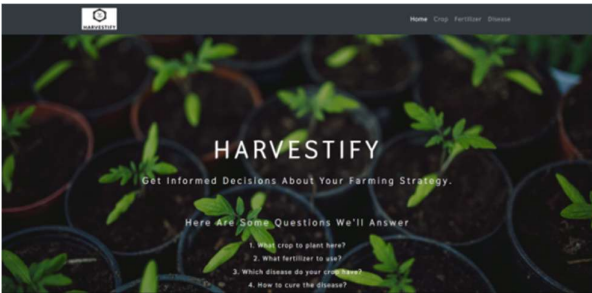


Figure 9

This screenshot is of a mockup of the "HARVESTIFY" website, which is a farming strategy application. The page has a dark, hazy background image of potted seedlings to give it a natural ambience. It guarantees well-informed decisions regarding planting crops, fertilizers, and pests, with a clean navigation bar at the top.

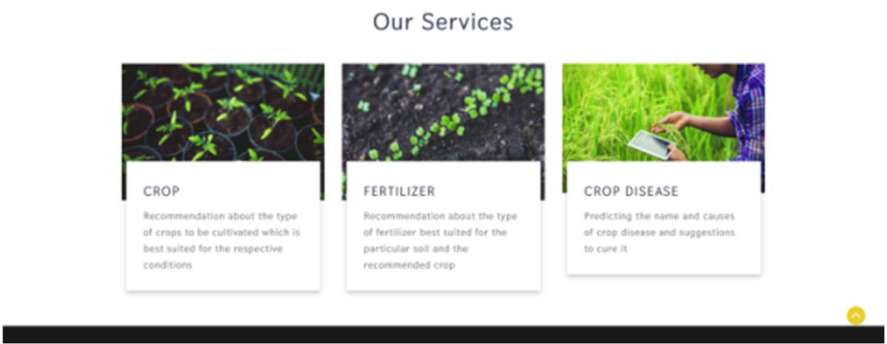


Figure 10

This picture shows a website page entitled "Our Services" and consists of three columns: "Crop," "Fertilizer," and "Crop Disease." There is an image and a short description of the service provided in each column, e.g., crop suggestion, selection of fertilizer, and prediction of disease. The design is minimalistic and contemporary.

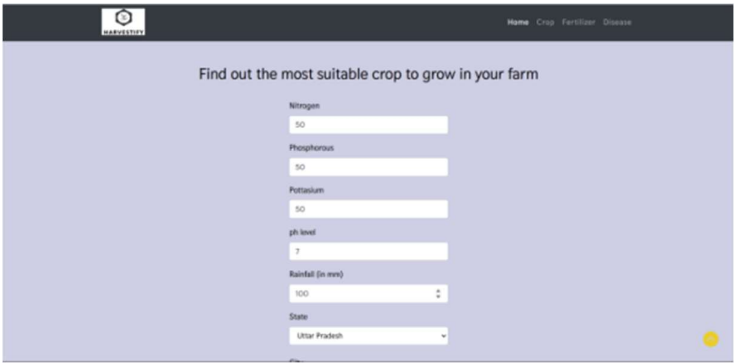


Figure 11

This image shows a webpage for "HARVESTIFY," designed to help users find the most suitable crop for their farm. It features input fields for soil parameters like Nitrogen, Phosphorus, Potassium, pH level, Rainfall (in mm), and State (pre-filled as Uttar Pradesh). A simple, clean interface is used.

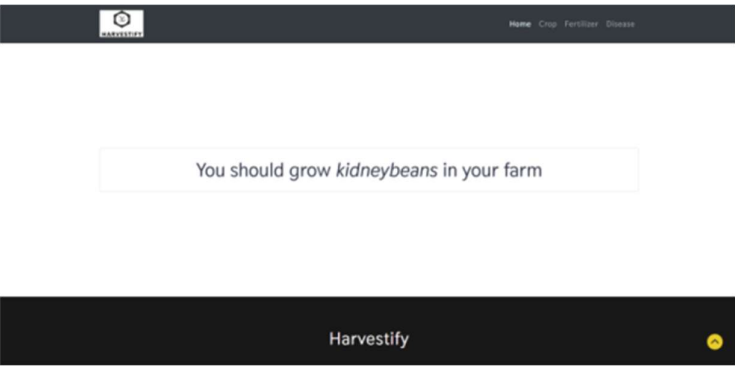


Figure 12

This is a screenshot of the results page of the "HARVESTIFY" website. It presents the suggestion "You should grow kidneybeans in your farm" in a central box. The page is minimalistic, with a white background, dark header and footer, and the "HARVESTIFY" logo at the top left corner.

Figure 13

Here's a screenshot of a page from "HARVESTIFY" for fertilizer guidance. Users provide soil values (Nitrogen, Phosphorus, Potassium) and a choice of target crop (apple). A "Predict" button triggers the prediction. The site is designed cleanly with a pale background and black header/footer.

Figure 14

This photo shows fertilizer suggestions from "HARVESTIFY." It reads "The K value of your soil is low" and proposes solutions such as applying potash fertilizers, kelp meal, Sui-Po-Mag, or planting banana peels. The suggestions are in a light-yellow, centered box on a white, clean webpage.

Figure 15

This is a screenshot of one page from "HARVESTIFY" employed for plant disease diagnosis. There is a photograph of a plant uploaded by the users; in this instance, there are rust-like spots on the leaves. The webpage contains "Find out which disease has been caught by your plant" and has a "Predict" button. It is minimalist with a light-colored background.

Figure 16

This photo shows a disease diagnosis outcome from "HARVESTIFY." It labels "Bacterial Spot" on a "Pepper" plant. The pathogen is outlined in detail, listing the *Xanthomonas* bacteria involved and how to see them under a microscope. It's set in a centered box on a white, clean-looking webpage.

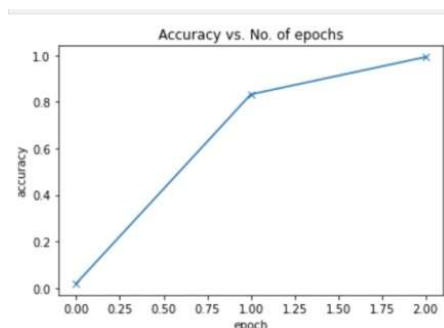


Figure 17

This is a line graph labeled "Accuracy vs. No. of epochs." It represents a model's accuracy rising with training epochs. The graph is low, takes a sharp rise to a point of 0.8 accuracy at epoch 1, and then takes a more gradual rise to 1.0 at epoch 2.

V. CONCLUSIONS

A. Performance Evaluation

The regression models that were developed within this project proved to be most accurate in assessing and predicting academic performance of the students. Attending classes, study time, and socio-economic status were some of the identified predictors of success with the highest magnitude. The performance evaluation measures, i.e., R-squared, Mean Absolute Error (MAE), and Root Mean Squared Error (RMSE), of the performance set up the reliability and robustness of the models. The comparison between actual and forecast score showed a high congruence, ensuring that the chosen method and instruments are appropriate. That kind of finding is enough evidence to assert that regression models have the capability of effectively measuring and capturing the influence of various factors and the performance of students. The analysis also found subtleties about the interaction between predictors. For example, while attendance and study time were positively related to academic attainment in a linear fashion, they had a more powerful effect when combined with parental levels of education. The results provide insight into the multifaceted manner in which single and combined factors contribute to student achievement and provide teachers with practical routes for driving improvement. Key Contributions

This project offers several valuable contributions:

- Data-Driven Insights: The findings provide a data-backed understanding of the critical factors influencing academic performance.
- Predictive Modeling: The regression models can serve as foundational tools for educational institutions to forecast outcomes and identify at-risk students.
- Actionable Recommendations: The results highlight areas for intervention, such as promoting consistent attendance and engaging parents in the educational process.
- Scalability: The methodology and models developed in this project can be adapted for use in various educational contexts, from schools to universities, and across different student demographics.

B. Limitations

While the project achieved its objectives, several limitations were identified:

- Data Diversity: The dataset primarily represented a specific group of students, limiting the generalizability of the findings to broader populations.
- Linear Assumptions: Regression models assume linear relationships between variables, which may oversimplify complex interactions in real-world scenarios.
- Unobserved Variables: Certain external factors, such as psychological well-being, motivation, and classroom dynamics, were not included in the dataset, potentially affecting the comprehensiveness of the analysis.

C. Future Directions

To build upon this work, future studies and initiatives can focus on the following:

- Incorporating Additional Variables: Including psychological, behavioral, and environmental factors in the dataset can enhance the predictive power and provide a more holistic view of student performance.
- Exploring Non-Linear Models: Advanced techniques such as polynomial regression, decision trees, or ensemble methods can capture non-linear relationships and interactions between predictors.
- Real-Time Predictive Systems: Developing systems that integrate real-time data (e.g., attendance records, grades, and behavioral reports) can help educators identify and address issues proactively.
- Personalized Interventions: Using predictive models to design customized interventions for individual students based on their unique needs and challenges.
- Cross-Institutional Studies: Expanding the scope of analysis to include data from multiple schools or universities to validate and generalize the findings across diverse educational settings.

D. Broader Implications

This project showcases the potential of datadriven approaches to revolutionizing education. With the help of regression analysis, teachers and schools can move past intuition-led strategies and evidence-informed practice. The findings of this research can be used for anything from curriculum planning to teacher training, with the consequence that educational resources are both utilized effectively and equitably. In addition, by noting at-risk students, these institutions can act early, reducing rates of dropouts and, overall, enhancing student performance.

REFERENCES

- [1] S. M. Metev and V. P. Veiko, *Laser Assisted Microtechnology*, 2nd ed., R. M. Osgood, Jr., Ed. Berlin, Germany: Springer-Verlag, 1998.
- [2] J. Breckling, Ed., *The Analysis of Directional Time Series: Applications to Wind Speed and Direction*, ser. Lecture Notes in Statistics. Berlin, Germany: Springer, 1989, vol. 61.
- [3] S. Zhang, C. Zhu, J. K. O. Sin, and P. K. T. Mok, "A novel ultrathin elevated channel low-temperature poly-Si TFT," *IEEE Electron Device Lett.*, vol. 20, pp. 569–571, Nov. 1999.
- [4] M. Wegmuller, J. P. von der Weid, P. Oberson, and N. Gisin, "High resolution fiber distributed measurements with coherent OFDR," in *Proc. ECOC'00*, 2000, paper 11.3.4, p. 109.
- [5] R. E. Sorace, V. S. Reinhardt, and S. A. Vaughn, "High-speed digital-to-RF converter," U.S. Patent 5 668 842, Sept. 16, 1997.
- [6] (2002) The IEEE website. [Online]. Available: <http://www.ieee.org/>
- [7] M. Shell. (2002) IEEEtran homepage on CTAN. [Online]. Available: <http://www.ctan.org/tex-archive/macros/latex/contrib/supported/IEEEtran/>
- [8] FLEXChip Signal Processor (MC68175/D), Motorola, 1996.
- [9] "PDCA12-70 data sheet," Opto Speed SA, Mezzovico, Switzerland.
- [10] A. Karnik, "Performance of TCP congestion control with rate feedback: TCP/ABR and rate adaptive TCP/IP," M. Eng. thesis, Indian Institute of Science, Bangalore, India, Jan. 1999.
- [11] J. Padhye, V. Firoiu, and D. Towsley, "A stochastic model of TCP Reno congestion avoidance and control," Univ. of Massachusetts, Amherst, MA, CMPSCI Tech. Rep. 99-02, 1999.
- [12] Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specification, IEEE Std. 802.11, 1997.
- [13] Samani, A., & Alappatt, U. (2019). QC-Automator: Deep learning-based automated quality control for diffusion MR images. *Frontiers in Neuroscience*, 13, 1456.
- [14] Reiner, M. (2023). Automating quality assurance for digital radiography. *Journal of the American College of Radiology*, 6, 486.
- [15] Sagan, V. M., Gitelson, A. A., & Rundquist, D. C. (2021). Data-driven artificial intelligence for calibration of hyperspectral big data. *IEEE Transactions on Geoscience and Remote Sensing*, 59(12), 9551-9566.
- [16] Singh, A. K., Kumar, N., & Singh, A. P. (2023). Plant disease detection using UAV and deep learning techniques. *Encyclopedia MDPI*.
- [17] Goel, P., & Yadav, R. K. (2021). Smart agriculture – Urgent need of the day in developing countries. *ResearchGate*.
- [18] Zhang, H., Chen, M., Wang, G., Sun, H., & Jing, F. (2020). An automated early-season method to map winter wheat using time-series Sentinel-2 data: A case study of Shandong, China. *Computers and Electronics in Agriculture*, 182, 106135.
- [19] Shivam Pathak, Shivam, Tarun Yadav, Ashish Gupta.
- [20] Al-Zubaidi, A., & Al-Wadaani, S. (2023). Blockchain for precision irrigation: Opportunities and challenges. *Journal of King Saud University - Computer and Information Sciences*, 35(1), 101-113