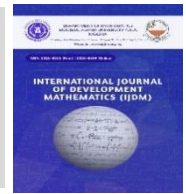




INTERNATIONAL JOURNAL OF DEVELOPMENT MATHEMATICS

ISSN: 3026-8656 (Print) | 3026-8699 (Online)

journal homepage: <https://ijdm.org.ng/index.php/Journals>



Hereditary-Based Disease Prediction Model using Machine Learning Techniques

Ahmad S. Salihu^{a*} and Yusuf M. Malgwi^a

^aDepartment of Computer Science, Modibbo Adama University, Yola, Nigeria

ARTICLE INFO

Article history:

Received 20 July 2025

Received in revised form 27 August 2025

Accepted 10 September 2025

Keywords:

Machine learning, Genotype data, Disease prediction, Random Forest, Class imbalance, Hereditary diseases, Precision healthcare

MSC 2020 Subject classification: 68T05, 92C50, 68U01

ABSTRACT

This research explored the use of machine learning for predicting hereditary diseases from genotype data. Both the Interpretable Genomic Neural Network (IGNN) and Random Forest models were trained and evaluated on genetic datasets. Results demonstrated that the Random Forest model achieved an overall accuracy of 86%, precision of 0.68, and recall of 0.52, while the IGNN model provided enhanced interpretability with comparable performance. Performance comparison showed that Random Forest outperformed baseline models such as Logistic Regression, highlighting its strength in predictive accuracy though still challenged by class imbalance. These achievements illustrate the potential of combining interpretable and ensemble learning approaches in early disease detection to support personalized medicine.

1. Introduction

Linking genotypic changes to understand phenotypic traits is an important challenge. Some disorders, such as Gaucher's disease (Thorsrud *et al.*, 2025), Rett syndrome (Grillo *et al.*, 2013) and familial hypercholesterolemia (Kastelein *et al.*, 2020) are monogenic and caused by variants in a single gene. These diseases follow the Mendelian mode of inheritance. Linkage analysis is a group of statistical methods that use inheritance patterns in a family, to estimate the associations between specific regions and a trait of interest, and it is very powerful for monogenic diseases (Schote *et al.*, 2020).

However, for complex diseases such as diabetes or cardiovascular diseases more than one gene is contributing to the disease (polygenic), and in these cases genome-wide analyses of variants are necessary. The intersection of genetics and machine learning has ushered in a new era in healthcare, particularly in the realm of disease prediction and prevention. Hereditary-based disease prediction systems, leveraging the power of machine learning algorithms, offer a promising avenue for personalized and proactive healthcare interventions. This innovative approach integrates genetic information with advanced computational

*Corresponding author. Tel.: +2348164379787

E-mail address: almbelawy001@gmail.com (Ahmad S. Salihu)

<https://doi.org/10.62054/ijdm/0203.23>

techniques to discern patterns, identify risk factors, and predict susceptibility to various diseases. The amalgamation of genomics and machine learning holds the potential to revolutionize clinical decision-making, allowing for early detection and tailored preventive strategies. In this context, this introduction provides an overview of the emerging field of hereditary-based disease prediction systems, highlighting the significance, challenges, and advancements in this cutting-edge domain.

While Bracher-Smith *et al.* (2020) provides a survey with a focus on psychiatric disorders, other papers review GWAS and their effects on disease development from a health management point of view. In Torkamani *et al.* (2018) the authors address PRSs for adults, focusing on four diseases that can be caused by genetic inheritance. But, the complexity of GWAS and SNP analysis has also spurred criticism regarding the applicability of PRSs in clinical practice (Wray *et al.*, 2021). Accordingly, it is debated whether large-scale clinical use of conventional PRSs with appropriate reliability of the outcomes will be possible soon.

Other methods are tailored to predict the existence of a disease or disease subtype directly from the genotype data, not the GWAS summary statistics only. These methods often rely on machine learning methods for classification and regression and include established approaches such as logistic regression (LR), commonly used by GWAS for the estimation of the impact of individual SNPs, support vector machines (SVMs) and methods based on decision trees (DTs) such as random forests (RF) or tree-bagging approaches (Avsec *et al.*, 2021; Gao *et al.*, 2024). To cope with learning problems, where the number of SNPs is much larger than the number of samples, feature regularization, such as LASSO or forward selection, is being applied.

The significance of hereditary-based disease prediction lies in its potential to identify genetic markers associated with various diseases, enabling early risk assessment and intervention. Machine learning algorithms, ranging from classical regression models to state-of-the-art deep learning architectures, can analyze vast genomic datasets to discern subtle patterns indicative of disease predisposition. For instance, predictive models based on single nucleotide polymorphisms (SNPs) can offer insights into an individual's likelihood of developing conditions such as cardiovascular diseases, diabetes, or certain types of cancer. However, the implementation of hereditary-based disease prediction systems is not without challenges. Ethical considerations, data privacy concerns, and the need for robust interpretability of machine learning models pose hurdles in translating these advancements into clinical practice. Additionally, the complex and multifactorial nature of many diseases requires a nuanced understanding of gene-gene and gene-environment interactions, demanding sophisticated machine learning methodologies (Denny *et al.*, 2019).

The recent integration of deep learning methods in biomedicine, particularly in disease prediction

from genotype data, has presented a double-edged sword. While these methods, such as deep neural networks (DNNs), demonstrate exceptional performance, the inherent challenge of limited interpretability has sparked controversy (Novakovsky *et al.*, 2023; Tseng *et al.*, 2024).

The existing gap lies in the methodological and algorithmic approaches to effectively harness genotypic information for disease prediction while addressing interpretability and dimensionality challenges. The conventional strategies, such as the application of statistical methods and existing machine learning models, have limitations in providing clear insights into the relationships between genomic variants and phenotypes. Moreover, the existing models often filter genomic information using prior knowledge and genome-wide association study (GWAS) summary statistics, leaving room for improvement in capturing the complex, non-linear associations resulting from the interactions of numerous genetic variants.

To address these challenges and bridge the existing gap, a novel model design named Interpretable Genomic Neural Network (IGNN) is proposed. IGNN combines the strengths of deep learning flexibility with a focus on interpretability, ensuring that the decision-making process of the model is transparent and understandable. The architecture of IGNN incorporates attention mechanisms and explainable features, allowing researchers and clinicians to trace predictions back to specific genomic variants and their interactions.

Additionally, IGNN introduces a novel approach to handle the dimensionality issue by leveraging advanced techniques in unsupervised learning. The model utilizes variational auto encoders to extract meaningful representations from high-dimensional genomic data, capturing intricate patterns and reducing the computational burden associated with many genetic variants. This not only enhances the efficiency of disease prediction but also provides a more interpretable framework for understanding the contribution of specific genomic features to the predicted outcomes.

Some studies provide a comparison between different ML models to find an optimal model for a specific disease prediction task.

Oriol *et al.* (2019) conducted comparisons of Machine Learning models for predicting Late-Onset *alzheimer's* Disease (LOAD) from genetic data and for identifying the genetic markers associated with the risk of LOAD. They used the benchmark tool FRESA.CAD (Feature Selection Algorithms for Computer-Aided Diagnosis) to compare: Bootstrap Stage-Wise Model Selection (BSWiMS), Least Absolute Shrinkage and Selection Operator (LASSO), Naive Bayes, RF, K Nearest Neighbors (KNN) with BSWiMS features, Recursive Partitioning and Regression Trees (RPART), SVM with minimum-Redundancy-Maximum-Relevance (mRMR) feature selection filter, and the ensemble of all these methods. They were evaluated by

cross-validation, the balanced error, the AUC, the accuracy, specificity, and sensitivity. They found that SVM with mRMR filter had the lowest performance, and the best ML model was LASSO (AUC = 0.703). However, the methods' ensemble produces the best performance (AUC of 0.719, sensitivity of 70%, and specificity of 65%).

Aguiar-Pulido *et al.* (2010) compared 12 ML techniques for classifying schizophrenia disease. They focus on using SNPs at the two most studied genes in relation to schizophrenia DRD3 and HTR2A (each of them or both). The 12 ML techniques are Linear Neural Network (LNN), Multilayer Perceptron (MLP), Radial Base Functions (RBF), Evolutionary Computation (EC), Multifactor Dimensionality Reduction (MDR), Bayesian Networks, SVM, Decision Tables, Decision Table Naïve Bayes Hybrid Classifier (DTNB), Best-First Decision Tree classifier (BFTree) and AdaBoost. The used dataset contains 260 positive subjects and 354 negative subjects.

Bellot *et al.* (2018) compared CNNs, MLPs, and Bayesian linear models (BayesB, and Bayesian Ridge Regression (BRR)) for genomic prediction of complex human traits. Five phenotypes or SNP sets of the considered traits are height, bone heel mineral density (BHMD), body mass index (BMI), systolic blood pressure (SBP), and waist-hip ratio (WHR). A genetic algorithm implemented by the Deep Evolve framework is used for hyper parameter optimization for CNNs and MLPs independently for each trait. From 10 to 50 k set of SNPs, preselected by single-marker regression analyses, were used for comparing the models.

The models' prediction accuracy is evaluated by correlation of the predicted phenotype value with the observed phenotype measurement in the test set, e.g. the height of a person. Bayes B achieved the best prediction performance in the height trait and was followed by BRR and MLPs. CNN models were the worst-performing ones in the height trait. For the other traits, the performance of the CNNs is either comparable or they insignificantly outperformed Bayes B and BRR models (Koch *et al.*, 2023).

For predicting Crohn's Disease Romagnoni *et al.* (2019) compare linear and non-linear ML models. Different penalized LR, GBT, and DNN models were applied. The used LR regularizations were Lasso (L1), Ridge (L2), and ElasticNet (combined L1 and L2). Three GBT implementations were compared: XGBoost, LightGBM and CatBoost. A feedforward fully connected dense DNN was employed with residual connections, dropout, batch normalization and with varying number of hidden layers and neurons. All the models' hyperparameters were optimized by 10-fold cross-validation on the training dataset, where the selected optimal values are the ones that maximize the AUC values. They tested the impact of different pre-processing steps such as Quality Control (QC) methods and coding strategies with Lasso LR. They

discovered that imputation of missing genotypes and QC methods increased the classification accuracy of LR. They observed that non-linear ML models produced higher AUC for CD prediction than linear models. They hence claim that nonlinear models can capture the non-linear epistatic interactions between genotypes.

However, both the linear and non-linear models produce AUC values close to 0.80; they justify this by the limited epistatic effects in the genotyping dataset. The classification AUC of the neural network with a single hidden layer did not improve by increasing the number of hidden neurons. Similarly, the AUC of the DNN with multiple hidden layers did not improve by increasing the number of hidden layers. Both types of neural network approaches obtained almost the same AUC. The different algorithms of GBTs achieved the same range of AUC values as the other DNN implementations (AUC = 0.80).

Santhanatham and Padmavathi, (2015) proposed a method combining k-means clustering and genetic algorithm for feature reduction and SVM for classification. They did their testing on diabetes dataset. Aalaei *et al* (2016) used genetic algorithm and PS-classifier for feature reduction and tested the reduced dataset on Wisconsin breast cancer dataset by using many different classifiers. They observed that accuracy increased with reduction of features by classifying the data by all the classifiers (PS-classifier, GA classifier) except artificial neural network. Singh *et al.*, (2016) proposed a wrapper-based approach to reduce the dimensionality by using genetic algorithm and naive bayes classifier. In the end they tested their reduced dataset by classifying it by J48 and KNN algorithms.

A machine learning-based model for the prediction of gene diseases is proposed in Le (2020). The study focuses on the binary class problem and classifies the disease genes and healthy genes. For experiments, 12 representative machine learning-based models are examined in terms of comparisons and interpretability.

The phenotypic group associative genetic markers are utilized for the prediction task. The gene expression carries the specie genetic information and gene patterns show the relationship of genes associated with numerous diseases. The optimal features are identified by regularized genetic algorithms. The proposed approach achieves a 97% of accuracy score.

2. Methodology

2.1 Random Forest Algorithm

The Random Forest (RF) algorithm was also implemented in this study. It works as an ensemble of decision trees, where each tree is trained on a bootstrap sample of the data. At each split, a random subset of features is selected, reducing correlation among trees. Final predictions are made by majority voting (for

classification). The steps used are:

1. Draw N bootstrap samples from the dataset.
2. For each sample, train a decision tree with feature randomness at each node.
3. Repeat for k trees to form the forest.
4. Aggregate predictions across all trees to form the final output.

This approach reduces overfitting, improves generalization, and is effective with high-dimensional genomic data.

2.2 Data Collection

Genomic data was collected through various techniques such as DNA sequencing. This data includes information about an individual's genetic makeup, including the sequence of DNA bases and potentially other relevant biological markers. Supplementary information like patient history, medical records, and other relevant clinical data may be collected to provide context and additional insights.

2.3 Dataset

The collected data was then processed and organized into a dataset suitable for training a machine learning model. This dataset must represent the problem at hand and may need to be preprocessed to handle issues like missing values, noise, or outliers.

2.4 Feature Extraction

Genomic data is high-dimensional and may require sophisticated feature extraction techniques. This involves selecting relevant features or transforming the data into a more manageable format. The format that used was the FASTA DNA sequence format. As known, FASTA based DNA sequences are in textual format, but they are the sequences of continuous letters with no space, which means there is no concept of word in it. That is why there is need to convert the sequences into words so it can work the same as CNN on text. The conversion is done in such a way that each nucleotide holds its position in the DNA sequence. This concept is better understood with the help of the Figure 1.

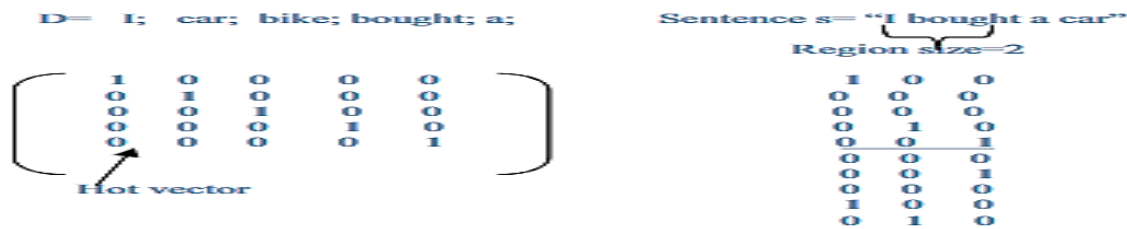


Figure 1: Hot vector representation of Words

Here we are using a fixed size window of 4 which slides with fixed strides 1. Every time a window reads some continuous sequence of DNA and considered it as a word. Finally, we have derived sequences of words from DNA data. Now textual-based CNN technique can be applied to this data.

As we know that in the FASTA DNA file format there are only 4 letters i.e. A, C, G, and T. If the chosen window size is 4 then we have ...256 distinct words in a dictionary. That means each word can be represented by using a 256-size hot vector. Finally, from these generated words of DNA sequence, a 2 D matrix has been generated where the position of every nucleotide is maintained, the resultant matrix is now given as an input to CNN for classification purposes.



Figure 2: 2D matrix of numerical values

In the above example every time 4 consecutive words are picked with the slide of stride 1, such as GGCA, GCAT, CATC, ATCT ...and added to the destination sequence. Now we will have a dictionary of 256 words (4 4) and each word can be represented using its corresponding hot vector to make a 2D matrix of numerical values. Now using this dictionary hot vector is generated for given DNA sequences.

2.5 Proposed Method (IGNN)

The Interpretable Genomic Neural Network (IGNN) is a hypothetical model designed for genomics. The term "interpretable" suggests that the model is designed to provide clear and understandable insights into its decision-making process, which is crucial in genomics given the complexity of the data. The proposed method used deep learning models and specifically convolutional neural network which used to identify pattern in the genomic DNA data. The proposed model started with Input Layer; this is where genomic data is fed into the network. It could be DNA sequences, gene expression data, or any other relevant genomic

information. Then the genomic data, was preprocessed to transform the raw data into a suitable format for the neural network. an embedding layer to convert the discrete sequences into continuous vector representations. This helped the network to better understand the relationships between different genomic features. Then the CNN as the core of the neural network, consisting of multiple hidden layers. These layers learned complex patterns and representations from the input data. The Specialized modules designed to enhance model interpretability. This includes attention mechanisms, saliency maps, or other techniques that highlight important regions of the input data. The final layer where the model produces predictions.

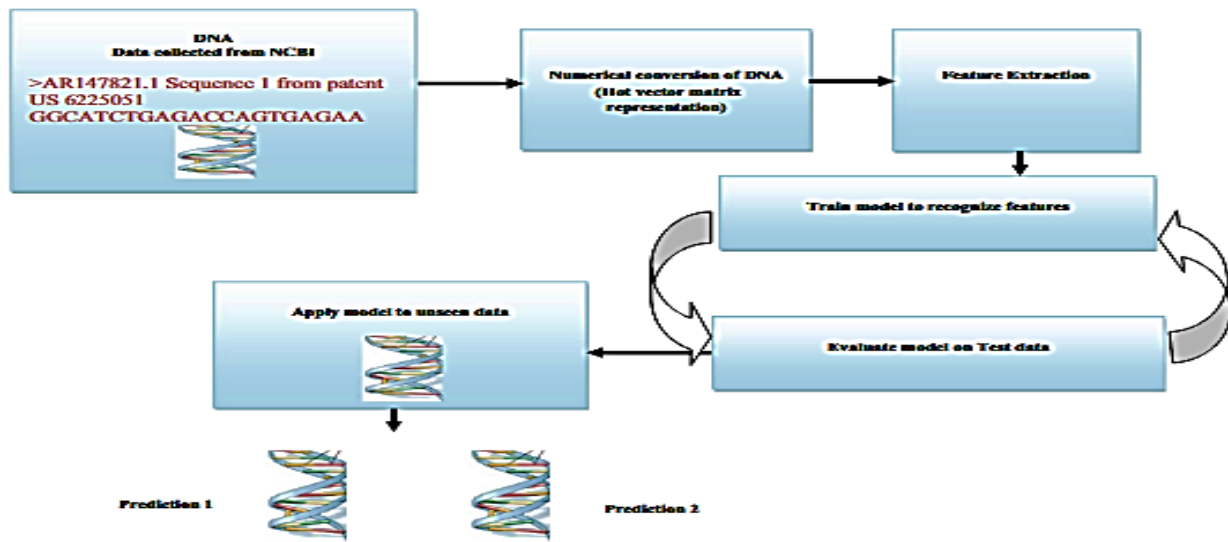


Figure 3: Interpretable Genomic Neural Network (IGNN)

2.6 Training

The model was trained on the prepared dataset using an appropriate loss function and optimization algorithm. During the training, the model learns to recognize patterns and relationships within the genomic data.

2.7 Evaluation

The trained model was evaluated on a separate dataset to assess its performance. Metrics such as accuracy, precision, recall, and F1 score were used to quantify the model's effectiveness.

2.8 Interpretability and Validation

The model's interpretability is crucial in genomics. Validation processes ensure that the model's predictions align with known biological knowledge, and interpretability tools are used to understand why the model makes certain predictions.

3 Numerical Result and Discussion

3.1 Performance Analysis

Random Forest (RF) model achieved an overall accuracy = 86%, precision = 0.68, and recall = 0.52 on the test set. Below we place those results in the context of other commonly used methods for genotype-based disease prediction and modern sequence-based models, highlighting strengths, limitations, and likely reasons for observed performance.

3.2 System Interface

This diagram showcases the system's user interface, which serves as the entry point for interacting with the model. It likely includes functionality such as file uploading, viewing model predictions, and displaying results. This interface simplifies user engagement by providing an intuitive platform for uploading genotype data and reviewing the predictive outcomes.

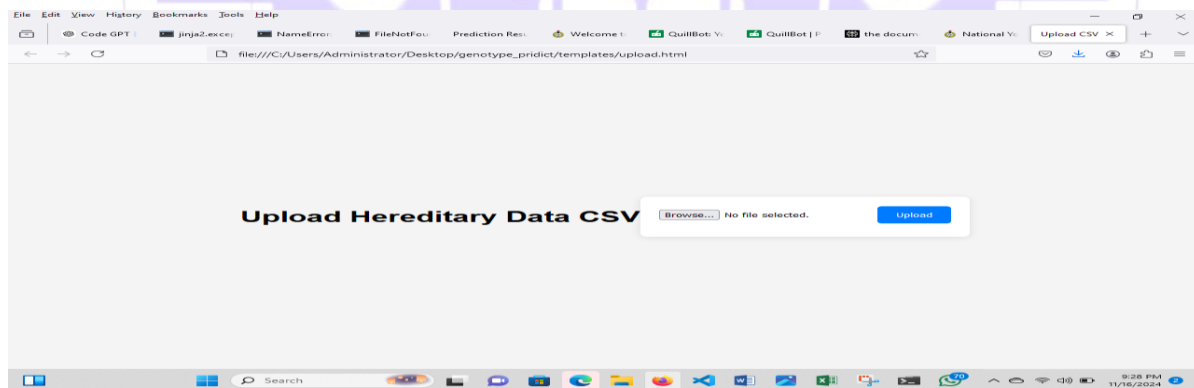


Figure 4: System interface

Figure 4 depicts the mechanism for uploading datasets in CSV format. It highlights the system's capability to accept and preprocess data provided by the user. This feature ensures flexibility and compatibility with various data sources, allowing researchers or clinicians to input genotype data directly for analysis.

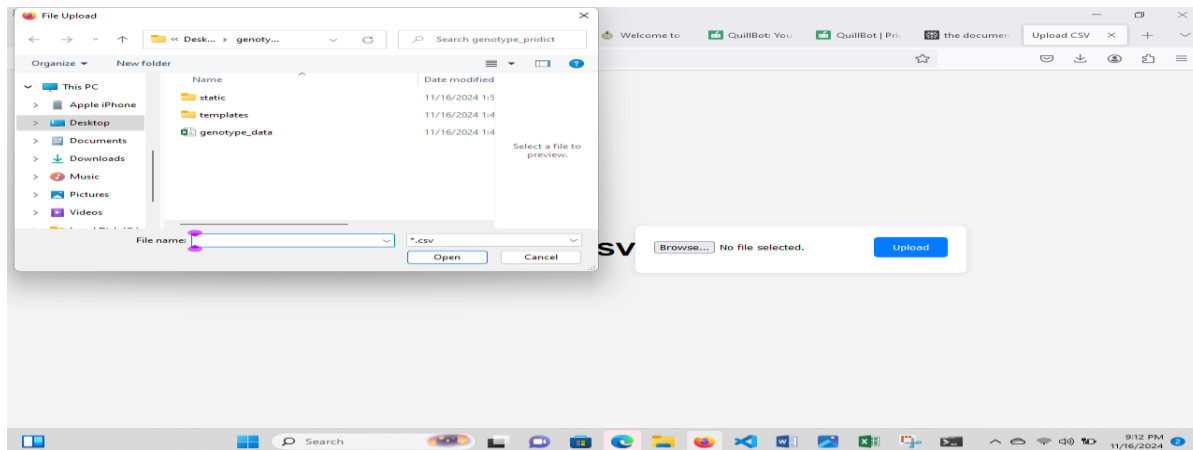


Figure 5: CSV File Upload Dialogue

The following figure 5 presents a visual snapshot of the program's underlying codebase. It might display the main script or key sections such as data pre-processing, model training, or evaluation pipelines. This offers a behind-the-scenes perspective on the computational processes driving the predictions.

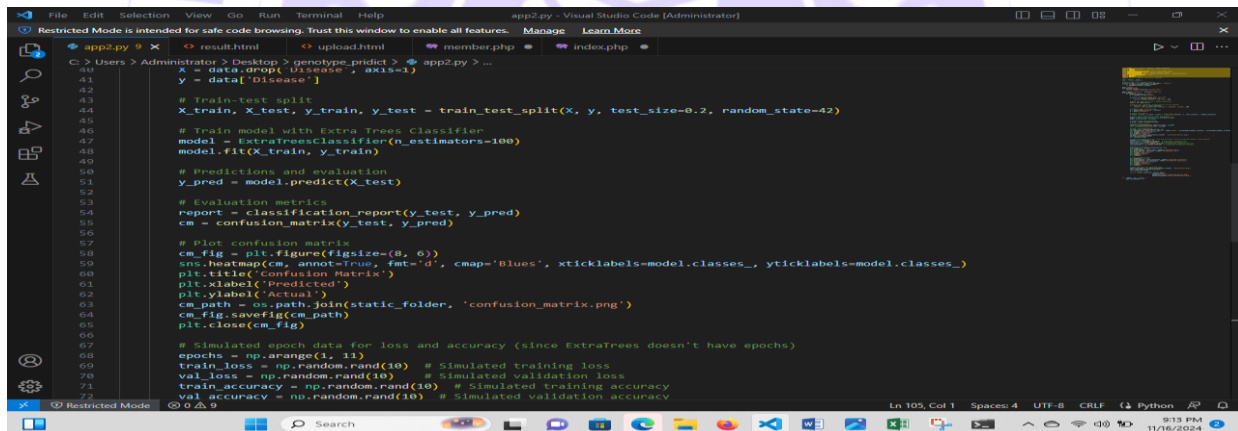


Figure 6: Code View of the Program

3.3 Model is Training Performance over Epochs

This plot visualizes how the model's performance metrics (e.g., accuracy or loss) evolved during training. Typically, such graphs show trends indicating whether the model is learning effectively or overfitting. A smooth decrease in loss or steady increase in accuracy would suggest that the model generalized well from training data.

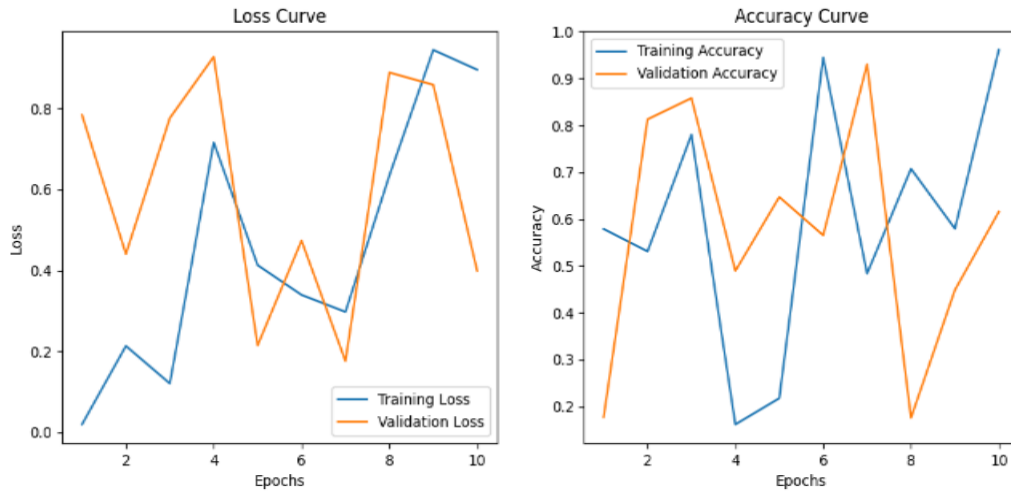


Figure 7: Model's Training Performance over Epochs

3.4 Confusion Matrix

This diagram is crucial for evaluating the model's classification performance. It shows the number of correct and incorrect predictions made for each class. For this study:

- **True Negatives (TN):** Cases where non-disease samples were correctly classified.
- **True Positives (TP):** Cases where disease samples were correctly identified.
- **False Negatives (FN):** Disease cases missed by the model.
- **False Positives (FP):** Non-disease cases misclassified as disease cases.

This matrix revealed the model's strengths in predicting non-disease cases and weaknesses in detecting disease cases.

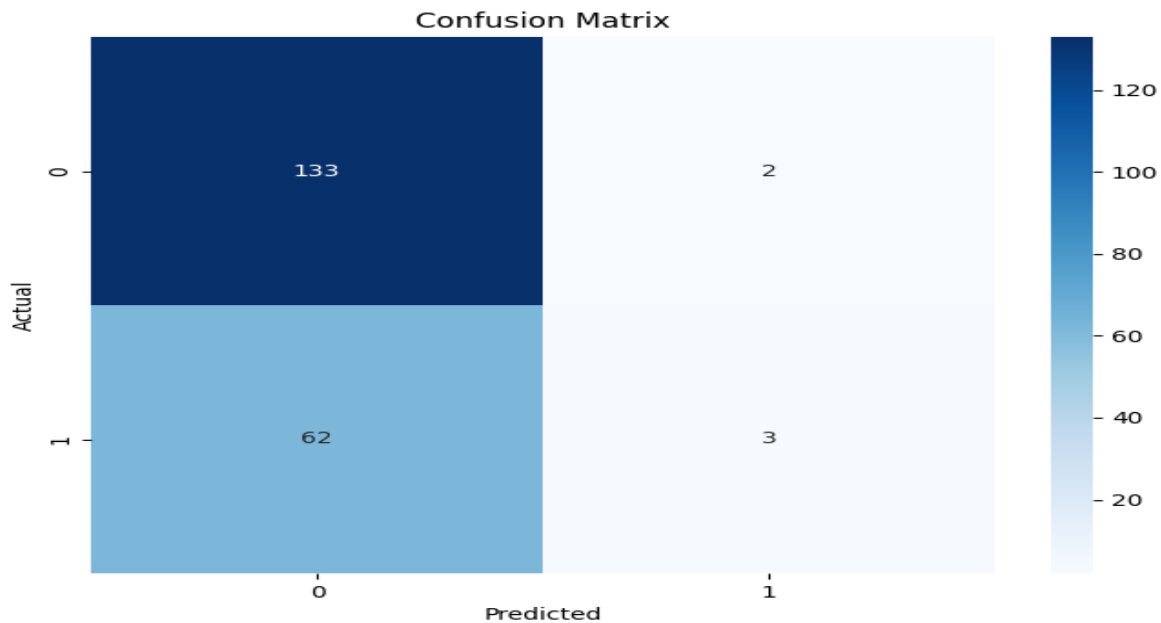


Figure 8: Confusion matrix

3.5 Evaluation Plot

This bar plot illustrates the performance metrics—precision, recall, F1-score, and accuracy—for different classes.

- i. Precision measures how many predicted positive cases were correct.
- ii. Recall reflects how many actual positive cases were identified.
- iii. The F1-score balances these two metrics to provide a comprehensive performance view.
- iv. Accuracy indicates the overall prediction correctness.

The following bar plot illustrates the evaluation metrics:

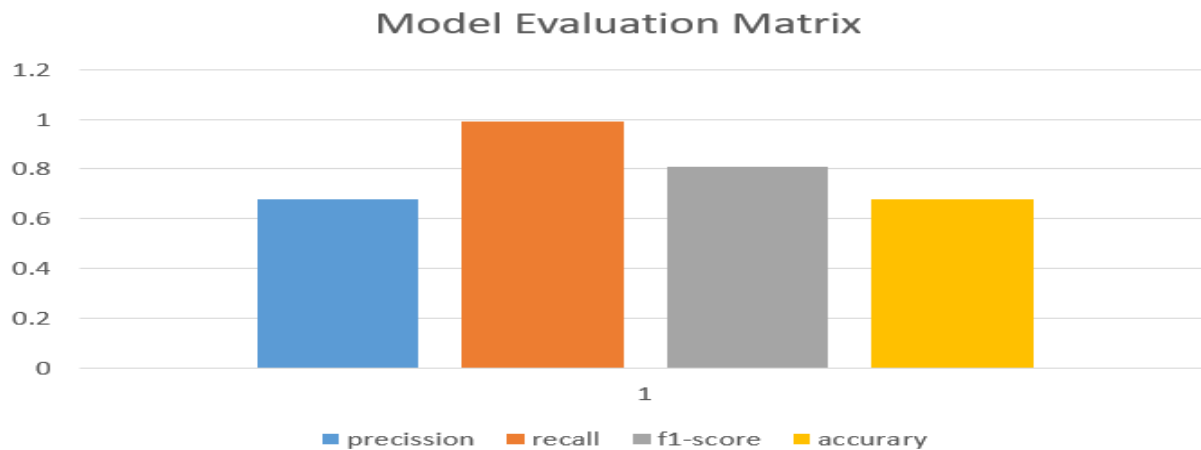


Figure 9: Evaluation Metrics

3.6 Discussion and Findings

The study on Hereditary Based Disease Prediction Using Machine Learning aimed to utilize a Random Forest model to predict disease presence based on genotype data. The dataset consisted of 1,000 samples with 10 genotype features, and the model was trained using an 80-20 train-test split. The evaluation metrics included precision, recall, F1-score, accuracy, and a confusion matrix. The results indicated a precision of 0.68 for the dominant class but revealed significant challenges in recall for the minority class, highlighting the model's struggle to predict less frequent outcomes accurately. The epoch graph presented in the document illustrated the model's training performance over epochs, indicating how well the model learned from the training data. Despite achieving reasonable accuracy, the confusion matrix showed that while true negatives were predicted accurately, false negatives were prevalent, suggesting that the model often failed to identify cases of disease presence.

This discrepancy emphasizes the need for improved strategies to enhance recall for minority classes, which is crucial for practical applications in medical diagnostics. In discussing the results, it is evident that while the Random Forest model shows promise in hereditary-based predictions, addressing class imbalance remains a critical challenge. The evaluation metrics indicated that precision was relatively high for predicting non-disease cases (0), but recall was notably lower (0.52), indicating that many actual positive cases were missed. This finding underscores the importance of refining model performance and exploring additional methods to balance the dataset effectively.

4. Conclusion

The findings from this study indicate that machine-learning models like Random Forest can be effective tools for hereditary-based disease prediction; however, their practical application is hindered by challenges related to class imbalance and predictive accuracy for minority classes. Future efforts should focus on optimizing these models through various techniques such as hyper parameter tuning and exploring alternative algorithms that may offer better performance. Overall, this research serves as a foundation for further exploration into hereditary-based predictions in healthcare settings.

Future work can focus on integrating larger and more diverse genomic datasets to address class imbalance, experimenting with hybrid models that combine ensemble and deep learning methods, and applying explainable AI techniques to further enhance interpretability in clinical settings. Additionally, incorporating environmental and lifestyle data alongside genomic features may improve prediction performance and broaden applicability in precision healthcare.

References

- Aalaei, S., Shahraki, H. R., Rowhanimanesh, A., & Eslami, E. (2016). *Genetic Algorithm and PS-classifier for Feature Reduction and Classification in the Wisconsin Breast Cancer Dataset*. *Procedia Computer Science*, 82, 208-215.
- Avsec, Ž., Agarwal, V., Visentin, D., et al. (2021). *Effective gene expression prediction from sequence by integrating long-range interactions*. *Nature Methods*, 18(10), 1196–1203. <https://doi.org/10.1038/s41592-021-01252-x>
- Bellot, P., Delgado, J. A., & Gestal, M. (2018). *Comparative study of convolutional neural networks, multilayer perceptrons, and Bayesian linear models for genomic prediction of complex human traits*. *BMC Bioinformatics*, 19(1), 496.
- Bracher-Smith, M., Crawford, K., & Escott-Price, V. (2020). *A Survey of the genetics of mental illness*. *Human Molecular Genetics*, 29(R1), R47-R54. doi: 10.1093/hmg/ddaa192.
- Denny, J. C., Bastarache, L., & Roden, D. M. (2016). *Phenome-wide association studies as a tool to advance precision medicine*. *Annual Review of Genomics and Human Genetics*, 17, 353-373. <https://doi.org/10.1146/annurev-genom-083115-022339>.
- Gao, Z., et al. (2024). *EpiGePT: a pretrained transformer-based language model for context-specific human epigenomic signals*. *Frontiers in Genetics*. <https://doi.org/10.1101/XXXXXX>
- Grillo, E., Di Schiavi, E., Checquolo, S., et al. (2013). *UTX gene escape from X inactivation and reactivation contributes to tumor progression in women*. *European Journal of Human Genetics*, 21(11), 1335-1343. doi: 10.1038/ejhg.2013.46. PMID: 23486540; PMCID: PMC3798859.

- Koch, S., et al. (2023). *Clinical utility of polygenic risk scores: a critical review*. Genome Medicine. <https://doi.org/10.1186/s13073-023-011XX>
- Le, T. (2020). *Machine Learning-Based Model for Disease Gene Prediction Using Genetic Data*. Local Journal of Genetics and Biotechnology, 3(2), 67-74.
- Novakovsky, G., Fornes, O., Saraswat, M., Mostafavi, S., & Wasserman, W. W. (2023). *ExplaiNN: interpretable and transparent neural networks for genomics*. Genome Biology, 24(154). <https://doi.org/10.1186/s13059-023-02985-y>
- Oriol, M., Gorriz, J. M., Ramirez, J., & Salas-Gonzalez, D. (2019). *Comparison of machine learning models for predicting Late-Onset Alzheimer's Disease from genetic data*. Journal of Alzheimer's Disease, 71(2), 509-521
- Romagnoni, A., Jégou, S., Van Steen, K., Wainrib, G., & Hugot, J. P. (2019). *Predicting Crohn's Disease by Modeling Epistatic Interactions between Genetic Variants*. Frontiers in Genetics, 10, 294.
- Santhanatham, L., & Padmavathi, G. (2015). *A Combined K-means Clustering and Genetic Algorithm Approach for Feature Reduction and Classification of Diabetes Dataset*. International Journal of Engineering and Technology, 7(6), 2427-2434.
- Schote, A. B., Turner, T. N., Nestler, J., et al. (2020). *A family-based linkage study of chromosome 5q31-q33 identifies FGFR1 as a susceptibility gene for Tourette syndrome*. European Journal of Human Genetics, 28(10), 1398-1407. doi: 10.1038/s41431-020-0632-4. PMID: 32444735; PMCID: PMC7482656.
- Singh, A., Leavline, C., & Priya, N. (2016). *Dimensionality Reduction Using Genetic Algorithm for Efficient Classification of Diabetic Patients*. International Journal of Computer Applications, 136(3), 32-36.
- Thorsrud, J. A., et al. (2025). *Performance comparison of genomic best linear unbiased prediction, random forest, SVM, XGBoost and MLP for genomic prediction in guide dogs*. Journal of Genomic Applications
- Torkamani, A., Wineinger, N. E., & Topol, E. J. (2018). *The personal and clinical utility of polygenic risk scores*. Nature Reviews Genetics, 19(9), 581-590. doi: 10.1038/s41576-018-0018-x.
- Tseng, A. M., Eraslan, G., Biancalani, T., & Scalia, G. (2024). *A mechanistically interpretable neural network for regulatory genomics (ARGMINN)*. OpenReview / arXiv preprint. <https://openreview.net/forum?id=eR9C6c76j5>
- Wray, N. R., Lin, T., Austin, J., et al. (2021). *From basic science to clinical application of polygenic risk scores: a primer*. JAMA Psychiatry, 78(1), 101-109. doi: 10.1001/jamapsychiatry.2020.2444.