

Fetal Guard ML: Advanced Predictive Insights for Enhanced Prenatal Care and Maternal-Infant Well-being

Suchitra Vittala Naik
Department of CSE
RNSIT, Bengaluru

Riya
Department of CSE
RNSIT, Bengaluru

Shilpa
Department of CSE
RNSIT, Bengaluru

S Mamatha Jajur
Assistant professor
Department of CSE
RNSIT, Bengaluru

Abstract—Cardiotocography (CTG) is an essential medical monitoring method used during pregnancy to evaluate the health of the fetus and provide information about high- and low-risk situations. The primary objective of this work is to create a machine learning model capable of predicting the fetus's health. The dataset used in this study is extensive documentation from CTG exams, providing a thorough summary of fetal health indicators. The dataset is utilized for both training and evaluating the model through machine learning techniques, which enables the model to absorb complex patterns in the dynamics of fetal health. The accuracy and precision of the machine learning model's performance are assessed, providing insight into how effectively it predicts fetal well-being. By providing healthcare practitioners with a useful tool for the early detection of possible problems and proactive treatment of fetal health during pregnancy, this research enhances the application of technology in prenatal care.

Index Terms—Cardiotocography (CTG), Logistic regression, Support vector machine (SVM), K-nearest neighbors (KNN), Random forest, Decision tree

I. INTRODUCTION

Maternity and parturition represent significant milestones in human life, with obstetrical science aiming to ensure the health and welfare of both mother and baby. However, fetal mortality within the womb remains a persistent challenge, often attributable to considerations such as chronic hypoxia, intrauterine growth retardation, maternal complications, congenital fetal malformations, and chromosomal abnormalities.

To address this challenge, technological advancements in antenatal fetal assessment have emerged as critical tools in reducing fetal mortality rates and improving overall perinatal outcomes. Non-invasive methods such as cell-free fetal DNA detection and ultrasound-based monitoring technologies like cardiotocography (CTG) have revolutionized fetal health assessment, offering insights into fetal well-being without invasive procedures.[1]

However, the implementation of these technological breakthroughs in medical care requires a thorough evaluation of their rationale and effectiveness. The criteria for prenatal fetal testing should prioritize tests that offer actionable information

beyond what is provided by clinical examination alone, ensuring that the benefits outweigh any risks involved with the testing procedures.

Effective fetal well-being assessment tests offer significant benefits, enabling healthcare providers to identify compromised fetuses early and implement timely measures to enhance perinatal outcomes. These efforts could span from simple measures such as maternal bed rest and continued fetal monitoring to more complex medical interventions or even termination of pregnancy in severe cases.

Machine Learning (ML) has become an encouraging strategy for enhancing fetal health assessment by analyzing vast datasets and recognizing fine aspects that might not be apparent to human observers. ML algorithms, such as K-Nearest Neighbours (KNN), Random Forest, and Support Vector Classification (SVC), have demonstrated potential in classifying fetal and maternal health outcomes based on CTG data, thereby assisting healthcare professionals in making informed decisions about prenatal care.

The integration of ML models into current healthcare systems holds promise for improving perinatal care outcomes by augmenting the capabilities of healthcare providers in assessing fetal health. By harnessing information from various origins, ML algorithms can enhance the accuracy and efficiency of fetal health assessment, ultimately producing more beneficial results for both mothers and their unborn babies.[2]

In summary, advancements in technology and the integration of machine learning methods into fetal health evaluation represent significant progress in obstetric care. Utilizing these innovations allows healthcare professionals to refine their assessment of fetal health, leading to better perinatal outcomes.

This enhanced capability not only supports more accurate monitoring of fetal well-being but also contributes to the overall safety and health of both the mother and the newborn, thereby promoting a more secure and effective approach to prenatal care.

II. MATERIALS

During this analysis, We created several categories of models to classify data into multiple categories. We used algorithms

like logistic regression, Mode, random forest, Support Vector machine and Decision Tree, KNN.

A. Dataset

We analyzed a dataset comprising 2126 fetal Cardiotocography (CTG) records. These records were processed automatically, and various diagnostic features were measured. Additionally, the CTGs were independently classified by three specialist obstetricians, who collectively agreed on a classification label designated for each record. The categorization approach employed in this work to group fetal stages consisted of three classes: PATHOLOGICAL, SUSPECT, and NORMAL. Hence, our experiments focused on three-class classification using this dataset.

B. Classification

Specifically, instances where the fetal_health variable had a value of 1.0 were relabeled as “NORMAL”, while instances with a value of 2.0 were relabeled as “SUSPECT”, and instances with a value of 3.0 were relabeled as “PATHOLOGICAL”. This transformation not only simplified the depiction of the fetus's condition but also facilitated a more intuitive understanding of the dataset's outcomes.

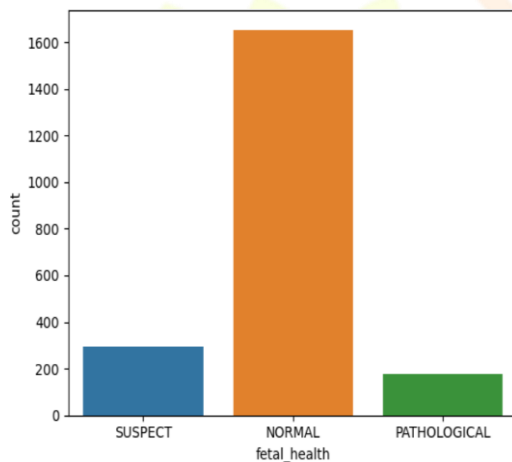


Fig. 1. Distribution of Cardiotocographic data

The CTG dataset exhibits a significant imbalance, comprising 2,126 samples in total: 1,655 are categorized as normal, 295 as suspect, and 176 as pathological, as depicted in Fig1

C. Correlation Matrix

Fig2 below acts as a key resource for examining the relationships between different variables in the dataset. The heatmap visualizes the correlation matrix, presenting a color-coded map of how variables interact with one another. This graphical representation helps to quickly spot trends and associations, highlighting areas of strong or weak correlation. By converting numerical values into a visual format, the heatmap makes it easier to understand complex data relationships, thereby facilitating more informed decision-making and insights.

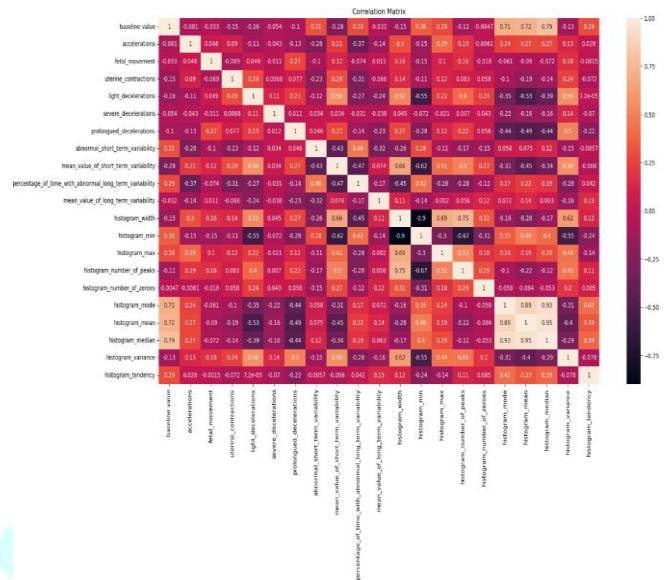


Fig. 2. Correlation Matrix

III. METHODOLOGY

Data pre-processing is a crucial initial step in data mining, turning raw data into reliable and usable information. It involves addressing missing values, outlier data, and inconsistent data. Missing values are handled through appropriate solutions, outliers are managed using specialized techniques, and inconsistent data is normalized for uniformity. The outcome is clean, standardized data ready for use in model development and evaluation. Typically, the dataset is split into subsets for model development and evaluation, with this study allocating 80% for model development and 20% for testing.

The model architecture comprises two distinct phases: data preprocessing and predictive modeling. Prior to these stages, data preprocessing is key to transform raw data into refined information suitable for subsequent analysis. This preliminary technique, integral to data mining, involves filtering and amalgamating data to establish a uniform format for further evaluation. Throughout the preprocessing phase, three primary issues are addressed: missing values, outlier data or noise, and inconsistent data. Missing values, resulting from errors or incomplete data collection processes, are rectified using appropriate missing data handling techniques. Outlier data, characterized by erroneous or irrelevant information, is mitigated through outlier data handling methods, targeting anomalies like mislabeling during data acquisition. Inconsistent data, arising from varied file formats or errors in data representation, is normalized to ensure uniformity and accuracy in subsequent analyses. This meticulous preprocessing stage lays the groundwork for robust and reliable predictive modeling, enabling effective data analysis and interpretation.[3]

1) *Splitting of Data*: The prediction model structure consists of two main components:

Training Model: This phase involves creating a training dataset with labeled features and targets. Algorithms applied are Mode, Logistic Regression, Support Vector Machine (SVM), Decision Tree, K-Nearest Neighbors (KNN), Random Forest, Gaussian Naive Bayes. These methods have been put to the test through experimentation on a variety of datasets that include categorical and numerical data. They are well-suited for classification tasks and can handle large datasets efficiently.[3]

Testing Model: During this stage, a testing dataset that lacks target labels but has the same properties as the training dataset is employed. The trained model predicts the target labels for the testing dataset based on its learned patterns and characteristics.

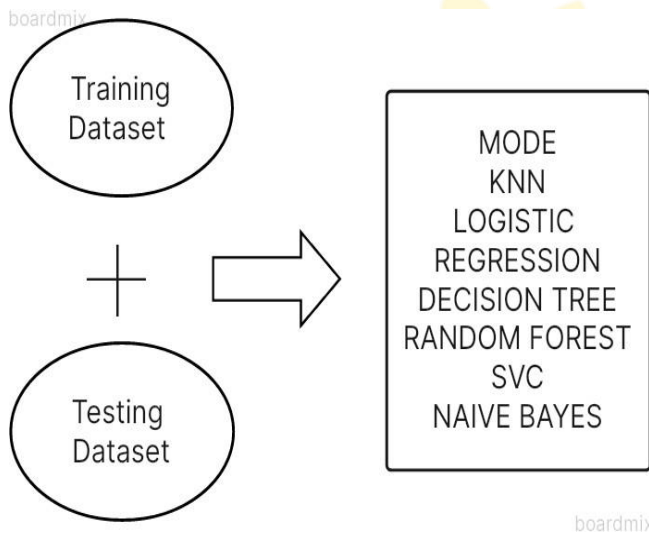


Fig. 3. Algorithm used for training and testing dataset

This structure ensures that the prediction model is trained on diverse algorithms and tested rigorously to produce reliable and accurate predictions for unseen data.

A. Algorithms

The listed algorithms are applied to train and test the dataset.

1) *Mode*: “Mode” refers to a fundamental statistical concept used to explain the central tendency of categorical data. It indicates the most commonly occurring value within a dataset or a particular feature. In situations where data points are classified into discrete groups or classes, like in classification tasks, this is very applicable. Understanding the mode helps analysts and practitioners gain insights into the distribution and prevalent characteristics of categorical variables, aiding in data exploration, preprocessing, and model evaluation.

2) *Logistic regression*: Logistic regression is a statistical approach employed for binary classification tasks in data analysis. Its purpose is to estimate the likelihood that a particular observation is identified as part of a class, utilizing one or more predictor variables. Unlike linear regression, which estimates continuous outcomes, logistic regression estimates probabilities through a logistic function, also referred to as the sigmoid

curve. This function ensures that the predicted probabilities fall between 0 and 1. It’s a powerful and interpretable tool for binary classification problems.[4]

3) *Support vector machine*: SVM employs a one-versus-rest approach for multiclass labeling, where each SVM predicts a specific class membership. The kernel choice transforms data points into linearly separable observations. Model optimization involves considering various parameters and employing search techniques to achieve the optimal combination. SVM utilizes an imaginary hyperplane to segregate classes and minimize structural risk, making it effective for small sample-learning. Kernel tricks enable SVM to handle non-linear data effectively.[4]

4) *KNN*: K-nearest neighbors (KNN) serves as an unsupervised learning algorithm designed to classify data by measuring the similarity among data points. It’s non-parametric and lazy, capable of categorizing data into multiple classes denoted by K. Using a majority vote among the K nearest neighbors, determined typically by Euclidean distance, KNN flexibly handles binary and multi-class classification tasks. Its simplicity and lack of explicit training make it ideal for scenarios with undefined data distributions or complex feature relationships.[5]

5) *Decision Tree*: Decision Trees are powerful models for both classification and regression tasks, with a preference for classification. They utilize a tree-like structure comprising decision and leaf nodes to influence decisions by considering dataset features. Algorithms like CART recursively partition datasets, constructing hierarchical structures that distill complex decision-making into intuitive paths. By posing binary questions and branching accordingly, Decision Trees offer valuable insights and predictive capabilities across diverse domains.

6) *Naive Bayes*: Naïve Bayes is a supervised learning algorithm rooted in Bayes’ theorem, commonly employed for classification tasks, especially in text classification with high-dimensional datasets. It’s renowned for its simplicity and effectiveness, facilitating rapid model building and quick predictions. As a likelihood-based classifier, Naïve Bayes predicts outcomes based on object probabilities, making it versatile in applications like spam filtration, sentiment analysis, and article classification. Its ease of use and applicability across various domains render Naïve Bayes a popular choice in machine learning workflows.

7) *Random Forest*: Random Forest is a popular algorithm for labeled Data Training Algorithm adept at handling both classification and regression tasks. It employs ensemble learning principles, integrating several tree-based classifier that have been trained using different dataset subsets to improve model performance. By aggregating predictions from individual trees using a majority voting mechanism, Random Forest mitigates overfitting risks and enhances predictive accuracy. This approach makes Random Forest a robust and widely utilized tool in machine learning.[4]

IV. RESULTS

A. Performance Assessment of Classification Models

1) **Accuracy:** The accuracy rate signifies the model's prediction precision, computed as the proportion of correctly classified samples by the classifier relative to the whole count of samples. This metric offers a clear evaluation of the model's effectiveness in making accurate predictions across the dataset.

$$\text{Accuracy Rate} = \frac{TP + TN}{TP + FN + TN + FP} \quad (1)$$

Here, TP represents True Positives, TN represents True Negatives, FN represents False Negatives, and FP represents False Positives.

2) **Recall rate:** The recall rate, Commonly called as true positive rate or sensitivity, represents the ratio of correctly predicted positive instances to all actually positive instances. Mathematically, it is defined as the frequency of true positives divided by the aggregate of correct positives and false negatives.

$$\text{Recall Rate} = \frac{TP}{TP + FN} \quad (2)$$

3) **Precision:** Precision quantifies the correctness of positive predictions by evaluating the proportion of correct positives among all instances identified as positive. It reflects the exactness of the predictive model. The formula for precision is given below:

$$\text{Precision} = \frac{TP}{TP + FP} \quad (3)$$

This equation represents the ratio of true positive predictions relative to the total quantity of cases classified as positive by the system. It provides useful information regarding the dependability and precision of the model's positive predictions.

B. Experimental Outcomes

The graph depicts the correctness or accuracy of various ML models. Logistic Regression has the highest accuracy, followed by SVC and Decision Tree. Logistic Regression exhibits the smallest difference between model training and performance evaluation accuracies, suggesting lower overfitting risk compared to Gaussian Naive Bayes, which shows the largest discrepancy. Additionally, as model complexity rises, accuracy decreases across all models due to increased likelihood of overfitting. All models show higher training accuracy compared to testing accuracy, indicating potential overfitting. Additionally, more complex models tend to have lower accuracy. Overall, Logistic Regression performs the best, but the most efficient model choice depends on specific data and goals.

The precision scores displayed in the graph Fig.5 represent both the model training data and performance testing data. Training records is employed to develop the model, while testing data remains separate and unseen during training, providing a more accurate indicator of the model's real-world performance. The scale of datasets can impact precision

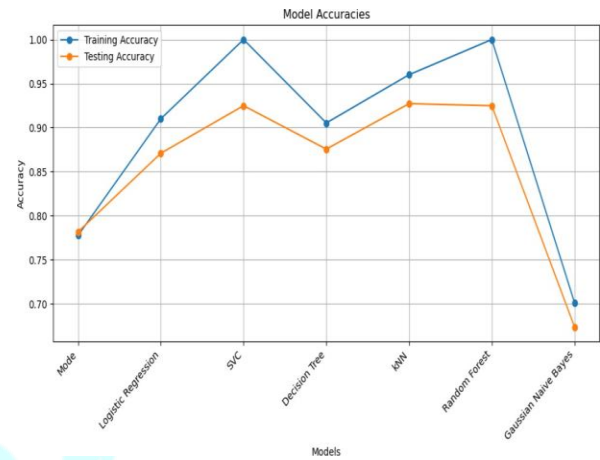


Fig. 4. Accuracy of various models

scores, with larger datasets often yielding higher precision. Additionally, model complexity influences precision, with more complex models potentially achieving higher scores but also risking overfitting to the learning data. Achieving a balance between model complexity and generalizability is fundamental to reliable predictions.

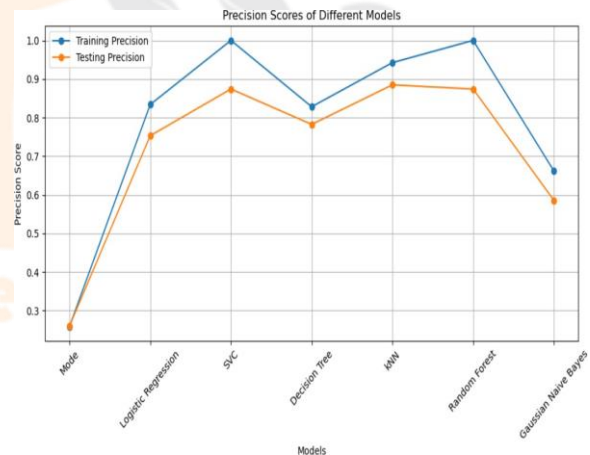


Fig. 5. Precision of various models

The graph Fig.6 illustrates that as scene complexity increases, model recall scores decrease across various scenes. Different models exhibit varying degrees of sensitivity to complex backgrounds, with some experiencing larger drops in recall scores than others. However, without defined parameters for scene complexity and size, the interpretation is somewhat limited, and results may vary according to the dataset used.

V. CONCLUSION

The paper delves into the critical importance of accurately predicting fetal health during pregnancy, emphasizing the complexities and likely challenges that could emerge, ultimately affecting the growth and well-being of the fetus. It underscores

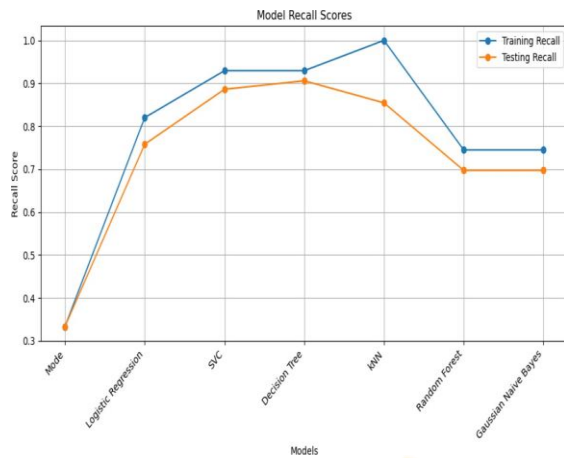


Fig. 6. Recall of various models

the requirement for proactive measures to recognize and mitigate risks before they escalate, thereby promoting optimal fetal development.

By leveraging various algorithmic models in machine learning, including Vector-Based Classification, Regression Analysis with Logistic Function, Random Forest Classifier, decision tree, and K-NN, the study aims to harness the power of data analysis to forecast fetal health outcomes. Through the evaluation of data obtained from cardiotocography (CTG), the research assesses the predictive capabilities of these algorithms, shedding light on their effectiveness in discerning patterns and trends indicative of fetal well-being or potential health issues.

The study highlights the significance of detecting fetal heart-beat deceleration as a vital sign of health status. Healthcare practitioners can use these findings to make timely and well-considered choices, intervening effectively when needed to ensure the health and survival of both mother and infant.

Furthermore, the analytical comparison of different algorithms underscores the superiority of Random Forest(RF) in the context of accuracy, highlighting its potential as a robust tool for improving clinical decision-making in obstetric care. This underscores the transformative impact of predictive modeling approaches in enhancing healthcare outcomes and reducing the incidence of adverse events during pregnancy and childbirth.

REFERENCES

- [1] Suhani Jain and Neema Acharya, "Fetal Wellbeing Monitoring: A Review Article" Cureus, vol. 14, no. 9, pp. e29039, Sep. 2022, doi: 10.7759/cureus.29039.
- [2] A. Mehbodniya, A.J.P. Lazar, J. Webber, D.K. Sharma, S. Jayagopalan, K. Kousalya, P. Singh, R. Rajan, S. Pandya, and S. Sengan, "Fetal health classification from cardiotocographic data using machine learning" Expert Systems, December 2021. DOI: 10.1111/exsy.12899.
- [3] Bahar, E., Agushinta R., D., Arimbi, Y. D., Reksoprodjo, M. (2022). "Model Structure of Fetal Health Status Prediction". JUITA: Jurnal Informatika, 10(2). e-ISSN: 2579-8901.
- [4] Chetan, Mohammed Rafi. (2020). "Fetal wellness Detection based on Machine learning techniques". International Research Journal of Engineering and Technology (IRJET), 07(09), e-ISSN: 2395-0056, p-ISSN: 2395-0072.
- [5] Md. Tamjid Rayhan, A S M Shamsul Arefin, and Sabbir Ahmed Chowdhury, "Automatic detection of fetal health status from cardiotocography data using machine learning algorithms" Journal of Bangladesh Academy of Sciences, vol. 45, no. 2, "pp. 155-167", January 2022.
- [6] Li, J., Liu, X. (2021). "Fetal Health Classification Based on Machine Learning". In Proceedings of the 2021 IEEE 2nd International Conference on Big Data, Artificial Intelligence and Internet of Things Engineering (ICBAIE). Zhuhai, China: IEEE. DOI: 10.1109/ICBAIE52039.2021.9389902.
- [7] Chinnaiyan, R., Alex, S., Vijayachinnaiyan, R. (2021). "Machine Learning Approaches for Early Diagnosis and Prediction of Fetal Abnormalities". In Proceedings of the 2021 International Conference on Computer Communication and Informatics (ICCCI -2021). Coimbatore, India: IEEE. DOI: 10.1109/ICCCI50826.2021.9402317.