

Machine Learning-Based Clinical Decision Support System for Predicting Fetal Hypoxemia during Non-Stress Test

Name:- Dr. Sajal Baxi

Guide:- Prof. Anoop Khanna

IIHMR- U/01/2019-21/0983

IIHMR University, Jaipur

MBA Health and Hospital Management

Title Slides	01
Introduction	03
Research Objectives	04
Literature Review	05-06
Methodology	07-10
Non-stress Test	11-13
Data and Data Structure	14-17
Algorithms and model	18-24
Model results	25-29
Discussion	30
Limitations	31
Conclusions	32
Reference Slides	33

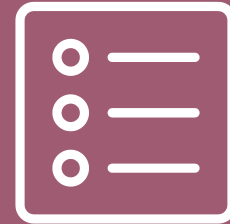


Table of contents

Introduction

01

Neonatal mortality rate refers to death of a live-born baby within the first 28 days of life. Early neonatal mortality refers to the death of a live-born baby within the first seven days of life, while late neonatal mortality refers to death after 7 days until before 28 days.

02

Neonatal mortality rate in India was 21.7 in 2020^[1] and remains the highest among developing countries according to UNICEF.

SDG 3 Goal 2 aims at reducing neonatal mortality rate to below 12 by 2030^[2].

03

Fetal hypoxia is the 3rd leading cause of neonatal mortality^[3].

Fetal Hypoxia is a condition in which the fetus does not receive adequate supply of oxygen. This can lead to CNS suppression and death in many cases. Other causes of neonatal mortality include Preterm birth, congenital disorders.

04

Neonatal mortality rate in rural areas is 1.5 times that of urban areas according to NFHS-4^[4].

Lack of proper healthcare systems and trained medical professionals is the cause of this disparity. AI and ML based Clinical decision support systems can help bridge the gap between demand and supply.

Research Objectives



Broad research question

How can Machine Learning based clinical decision support system be used to initiate an early response by identifying high risk pregnancies and prevent death due to fetal hypoxia?



Specific research objective - 1

To identify key features in the data that can be used to identify a pregnancy with high risk of developing fetal hypoxia.

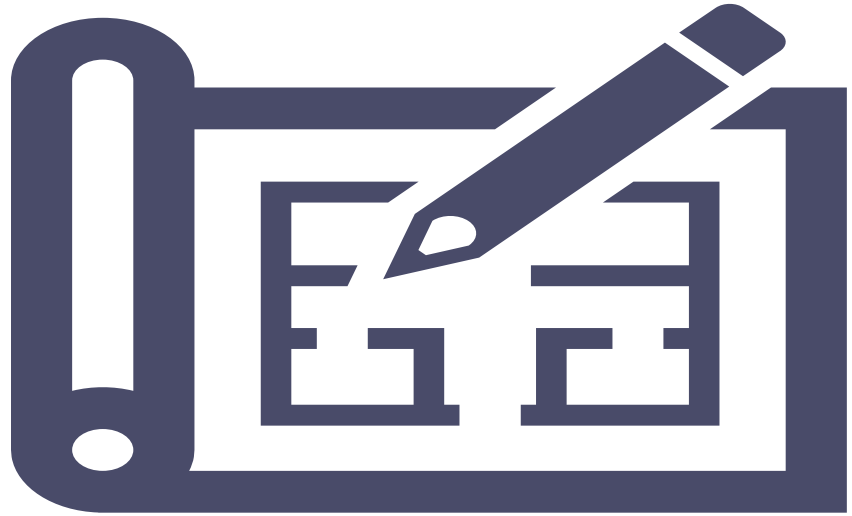


Specific research objective - 2

Using the features identified to develop a machine learning based predictive model.



Literature Reviews



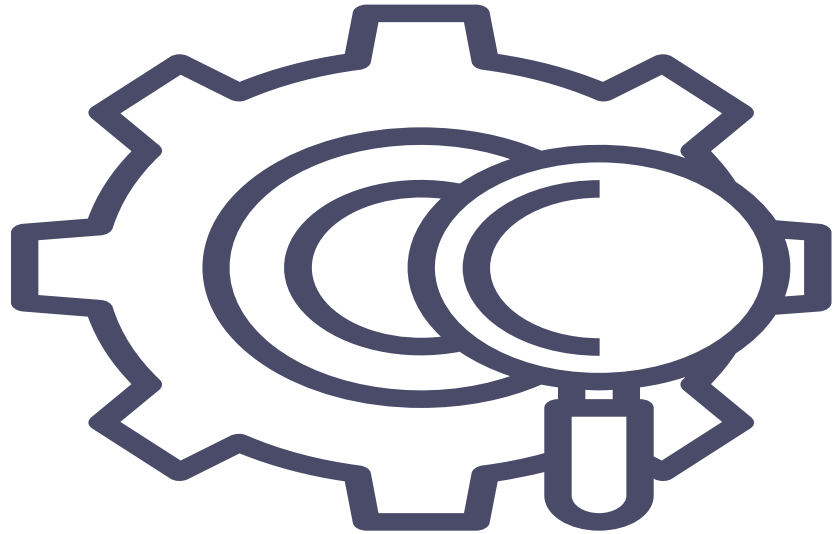
Literature review

Theory 01	Grivell, R., Alfircvic, Z., Gyte, G. and Devane, D. (2012). Antenatal cardiotocography for fetal assessment, Journal of Cochrane Database Syst Rev 1(2)	Hypoxia can be acute and chronic are two main types of fetal hypoxia while NST is applicable as monitoring test for fetus only after 7-month of pregnancy. Obstetricians choose caesarean section to avoid future negative outcomes after early detection of fetal hypoxia.
Theory 02	Ma, Q. and Zhang, L. (2015). Epigenetic programming of hypoxicischemic encephalopathy in response to fetal hypoxia, Journal of Progress in neurobiology 124: 28–48.	Maternal diseases, anemia during pregnancy, placental insufficiency and high-altitude pregnancy can be other causes of developing hypoxia.
Theory 03	Ayres-de Campos, D., Costa-Santos, C., Bernardes, J. and Group, S. M. V. S. (2005). Prediction of neonatal state by fetal heart rate tracings European Journal of Obstetrics Gynecology and Reproductive Biology 118(1): 52–60.	Non- stress test information is very useful for obstetricians in identifying future complications and to avoid permanent damage of fetus it must be diagnose early.

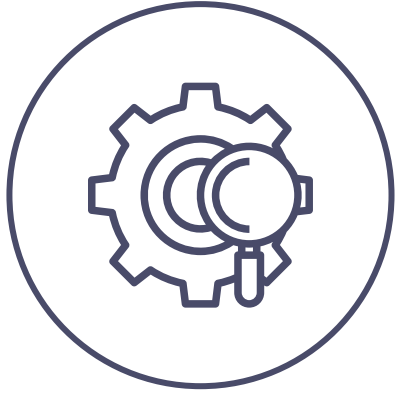
Literature review

Theory 04	Comert, Z., Kocamaz, A. and Gngr, S. (2016). Cardiotocography signals with artificial neural network and extreme learning machine, Proceedings of 24th IEEE Signal Processing and Communication Application Conference (SIU) 1542(9): 1493–1496.	International Federation of Gynecology and Obstetrics (FIGO) suggests some features like variability, number of acceleration and deacceleration patterns and baseline FHR are good indicators for fetal health.
Theory 05	Gribbin, C. and Thornton, J. (2006). Critical evaluation of fetal assessment methods. high-risk pregnancy management options, 3(2): 229–239.	FHR must have baseline between 110-160 beats per minute with variance of 5 or more for healthy fetus.
Theory 06	Dawson AJ, Buchan J, Duffield C, Homer CS, Wijewardena K. Task shifting and sharing in maternal and reproductive health in low-income countries: A narrative synthesis of current evidence. <i>Health Policy Plan.</i> 2014;29:396–408.	Using technological advancements, incorporation of an automated machine learning model with high accuracy (such as the one developed in this study) to predict suspect as well as pathological fetal state would enable task sharing with lay health-care providers to ensure timely referral and management of women in labor, hence improving perinatal outcomes.

Methodology



Methodology



Study design: Observational analytical



Methodology: CRISP – DM

It provides blueprint of the project by dividing it into six phases for better business understanding and comprehensive data mining.



Focus population: Females in 2nd or 3rd trimester of pregnancy.

Sampling method: Purposive sampling.

Sample size: 2126



The data is analyzed in R and models are build using various ML libraries like Boruta, Psych, RandomForest, ggPlot2.

CRISP- DM

Cross Industry Standard Process for Data Mining

Step1: Translating the prediction task of the research question into data mining process with models like classification, clustering. Research problem is a type of classification problem because it is based on 12 features from the clinical data like Fetal Heart Rate, Uterine contractions and predicting the Fetal health status which is a categorical feature.

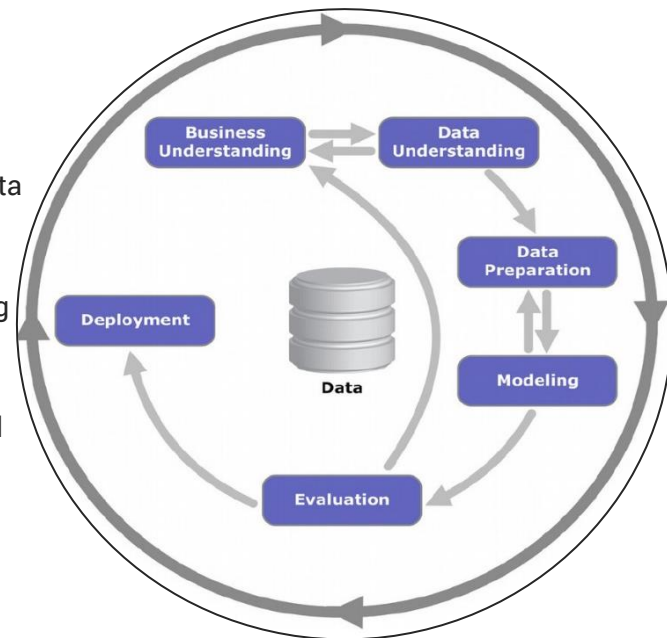
Step 2: Feature selection and developing insights from bar charts, flow models. Acquiring target attribute, finding relationships and statistical analysis are the main parts of initial data pre-processing.

Step 3: Improving the data quality by rectifying errors, transformation of values into existing variables.

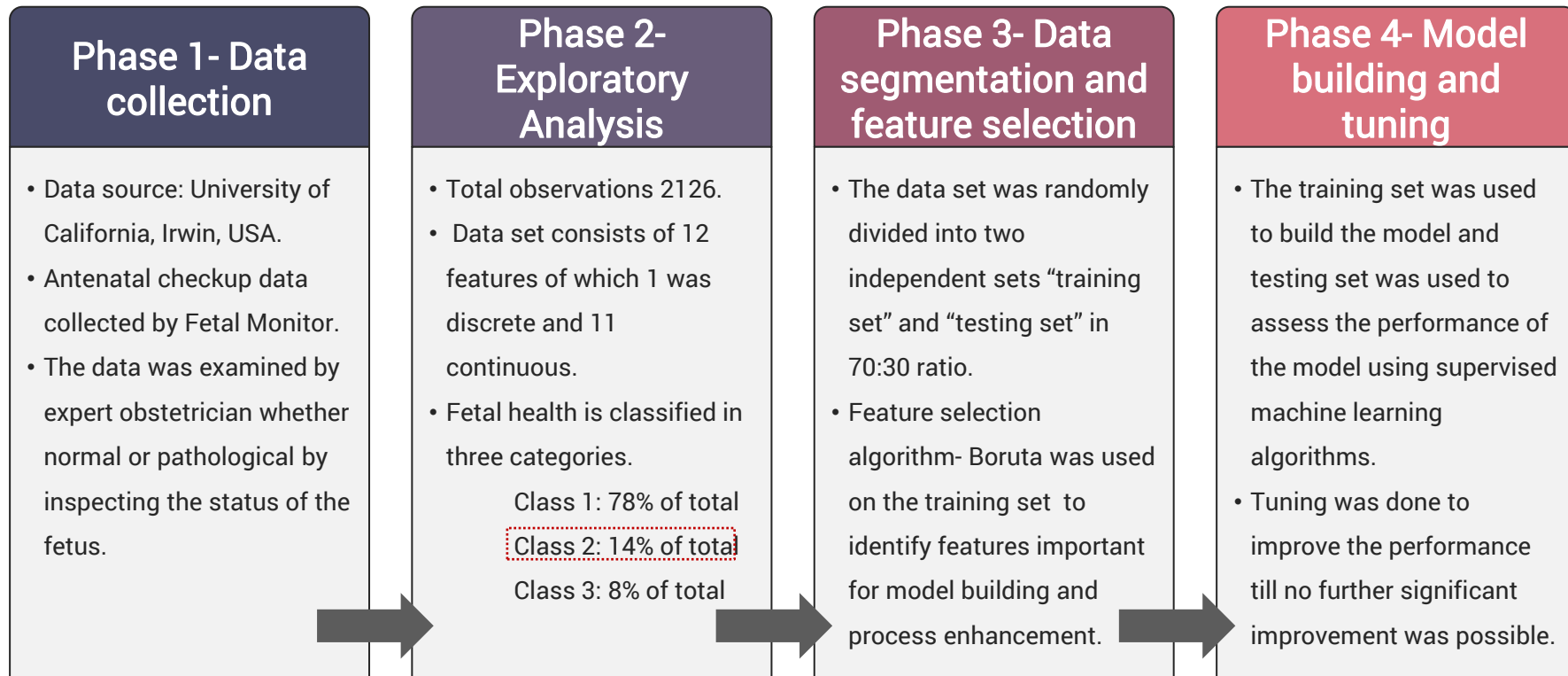
Step 4: Machine Learning algorithms like decision tree and RandomForest are used to build models due to higher accuracy and reduced overfitting properties of the algorithm.

Step 5: Accuracy, Sensitivity (Recall), Specificity, F-measure, Kappa Score are taken into consideration for prediction performance matrix in context of model evaluation.

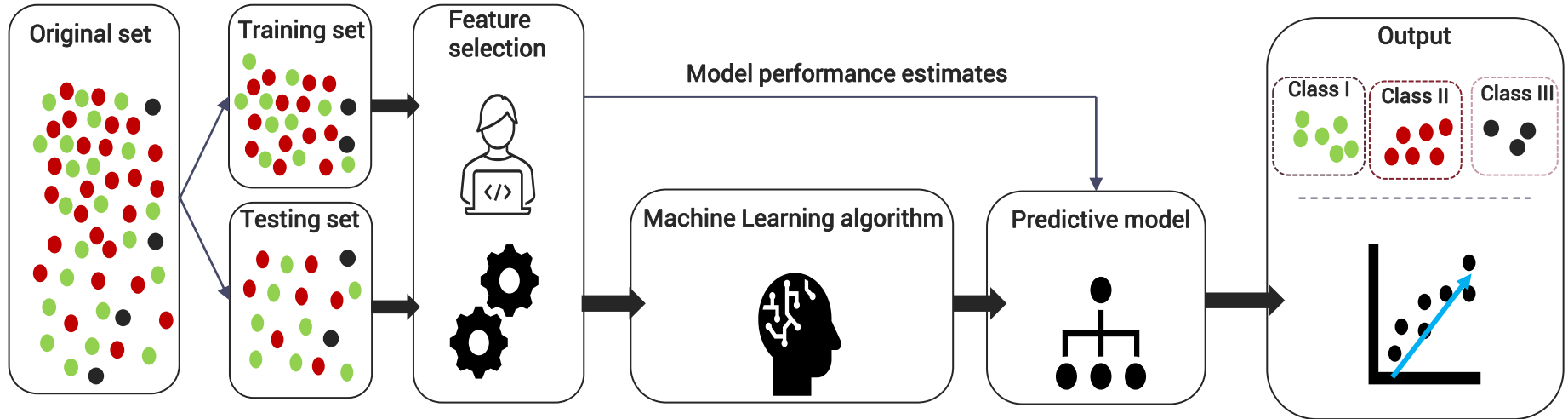
Step 6: Fine tuning of the model parameters to smoothen the processes.



Model building methodology



Supervised machine learning process



1. The real-world data may have unrecognized, missing and noisy elements in it. Thus, data cleaning and preparing is done as it helps in the model accuracy.
2. Data is then split into training and test data sets. The training set is used to train the model in the next step, while the test data is used to validate the model in the fourth step. The typical default is a 70/30 split between training and test sets.
3. The training set contains the predictor variable in it and is connected to the algorithms where sophisticated mathematical modeling is used to develop predictions.
4. Validation of the model is done using the testing set where the predictor variable is not known to the model. The predicted values are compared to the actual values to assess the model accuracy and other statistics.
5. The output obtained can either be discrete (divided into classes) or continuous based on the type of research problem.

Non stress Test

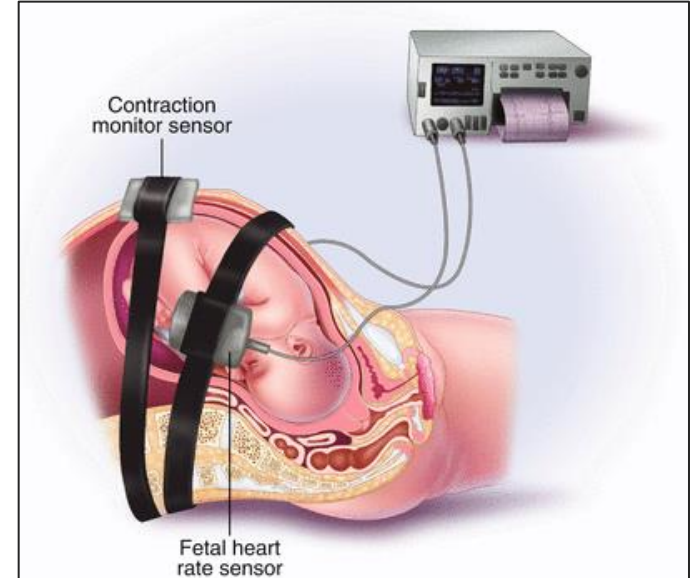
Non stress test(NST)^{[5][6]}

A nonstress test (NST) is a screening test used in pregnancy to assess fetal status by means of the fetal heart rate and its responsiveness.

1. It is a non-invasive procedure performed during routine antenatal checkup.
2. The test analyses fetal heart rate, its reaction to movements and uterine contractions.
3. It is done after 28th week of pregnancy when fetal neural reflexes are developed.
4. Fetal heart rate and accelerations are the most important recordings of the test as they reflect CNS alertness and activity.
5. The duration of the test is 20 minutes and if no significant readings are recorded the duration is extended to 60 minutes to remove the possibility of fetus being asleep.
6. Absence of accelerations are indicative of CNS depression due to hypoxia i.e., lack of oxygen supply to the brain.
7. The fetus is termed reassuring (reactive) or non-reassuring (non-reactive) based on the results obtained.
8. Interventions may range from intensive observation to an emergency cesarian section.

Procedure and data collection^[7]

1. The fetal heart rate and the activity of the uterine muscle are detected by two transducers placed on the mother's abdomen, with one above the fetal heart to monitor heart rate, and the other at the fundus of the uterus to measure frequency of contractions.
2. Pulsed ultrasound signals are emitted by a Doppler ultrasound transducer placed on the maternal abdomen and are reflected off the fetal heart back to the ultrasound transducer. Autocorrelation processing is then used to calculate fetal heart rate.
3. These transducer cables are connected to an external monitor that registers the readings on a special paper.
4. Registration of uterine contractions is done using an abdominal pressure transducer which converts the abdominal tension created by the contractions into a written signal, which provides information on both the frequency and the duration of contractions.



Interpretation of the test results

Based on the observations recorded during non-stress test the fetus is categorized into three type:

Reassuring. Non-reassuring and Abnormal features

	Reassuring Feature	Non-Reassuring	Abnormal Feature
FHR (bpm)	110-160	100-109	<100 or >180
Variability (bpm)	Up to 5	<5 (40 min to 90 min)	<5 (for more than 90 min)
Decelerations (15 per min)	No deceleration	Lasting up to 3 min	Lasting for more than 3 min
Uterine contractions (per hour)	Up to 5	More than 5	More than 10

Categorization of the values:

1. **Normal Trace:** Base line fetal heart rate in range of 110-160 beats per minute, variability of less than 5 beats per minute and no decelerations. i.e., all features fall in the reassuring category.
2. **Suspicious Trace:** Detection of three reassuring features and one non-reassuring feature.
3. **Pathological trace:** More than two non-reassuring features out of four detected.

Data and data structure

Data from non-stress test

Predictor variable

Target variable

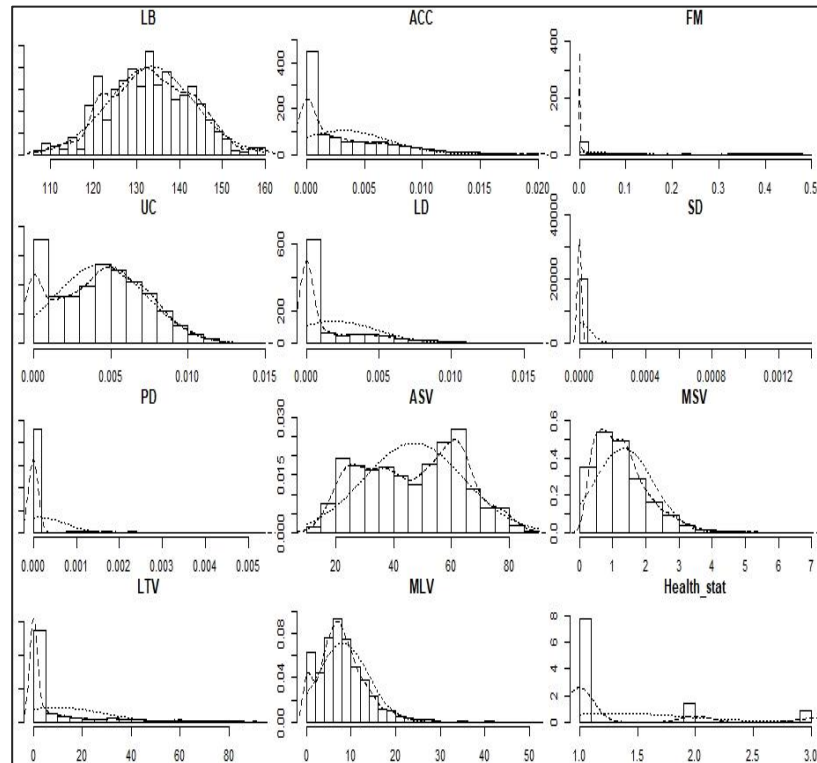
	A	B	C	D	E	F	G	H	I	J	K	L
1	LB	ACC	FM	UC	LD	SD	PD	ASV	MSV	LTV	MLV	Health_stat
2	120	0	0	0	0	0	0	73	0.5	43	2.4	2
3	132	0.00638	0	0.00638	0.00319	0	0	17	2.1	0	10.4	1
4	133	0.003322	0	0.008306	0.003322	0	0	16	2.1	0	13.4	1
5	134	0.002561	0	0.007682	0.002561	0	0	16	2.4	0	23	1
6	132	0.006515	0	0.008143	0	0	0	16	2.4	0	19.9	1
7	134	0.001049	0	0.010493	0.009444	0	0.002099	26	5.9	0	0	3
8	134	0.001403	0	0.012623	0.008415	0	0.002805	29	6.3	0	0	3
9	122	0	0	0	0	0	0	83	0.5	6	15.6	3
10	122	0	0	0.001517	0	0	0	84	0.5	5	13.6	3
11	122	0	0	0.002967	0	0	0	86	0.3	6	10.6	3
12	151	0	0	0.000834	0.000834	0	0	64	1.9	9	27.6	2
13	150	0	0	0.000983	0.000983	0	0	64	2	8	29.5	2
14	131	0.005076	0.072335	0.007614	0.002538	0	0	28	1.4	0	12.9	1
15	131	0.009077	0.22239	0.006051	0.001513	0	0	28	1.5	0	5.4	1
16	130	0.005838	0.40784	0.00417	0.005004	0	0.000834	21	2.3	0	7.9	1
17	130	0.005571	0.380223	0.004178	0.004178	0	0.001393	19	2.3	0	8.7	1
18	130	0.006088	0.4414	0.004566	0.004566	0	0	24	2.1	0	10.9	1
19	131	0.001524	0.382622	0.003049	0.004573	0	0.001524	18	2.4	0	13.9	2
20	130	0.002845	0.450925	0.00569	0.004267	0	0.001422	23	1.9	0	8.8	1
21	130	0.005055	0.46925	0.005055	0.004212	0	0.000842	29	1.7	0	7.8	1
22	129	0	0.340045	0.004474	0.002237	0	0.003356	30	2.1	0	8.5	3
23	128	0.004688	0.425	0.003125	0.003125	0	0.001563	26	1.7	0	6.7	1
24	128	0	0.333841	0.003049	0.003049	0	0.003049	34	2.5	0	4	3

Data Attribute Description

Attribute	Interpretation	Normal values
LB	Baseline Fetal heart rate (bpm)	110- 160 beats per minute
ACC	Accelerations per second	15 bpm increase for 15 seconds
FM	Fetal movements per second	10 per hour
UC	Uterine contractions per second	Up to 5 per hour
LD	Decelerations (Light) per second	15 bpm decrease for 15 seconds
SD	Decelerations (Severe) per second	15 bpm decrease for 20-30 seconds
PD	Decelerations (Prolonged) per second	15 bpm decrease for more than 90 seconds
LTV	Duration (%) of abnormal long-term variability	> 25 bpm
ASV	Duration (%) of short-term variability	3-5 bpm
MSV	Average short- term variability	Short term variability in 10-minute period
MLV	Average long-term variability	Long term variability in 10-minutes period
Health_stat	Fetal category (1= Normal; 2= Suspected; 3= Pathological)	Classification of the fetus

Data Structure

1. Data is collected from the fetal heart monitor readings on a per second basis over a period of 20-30 minutes.
2. The total sample size of the data is 2126.
3. Readings from the fetal monitor are obtained under 11 headers such as baseline heart rate, accelerations, number of fetal movements, uterine contractions, decelerations, Fetal health status.
4. Average short-term and long-term variability are derived continuous variables that are obtained from the fetal monitor.
5. Fetal health status is a derived discrete variable obtained after the obstetrician analysis the data obtained from the test.
6. Feature engineering is required to reduce complexity of data as well as to improve the model performance. Thus, all categorical variables are converted into numerical variables because it only accepts 1,2 or 3.

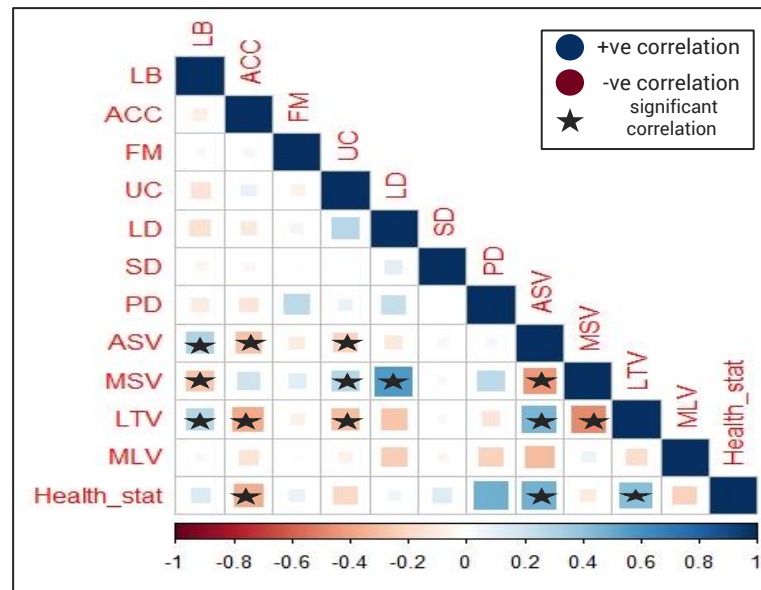
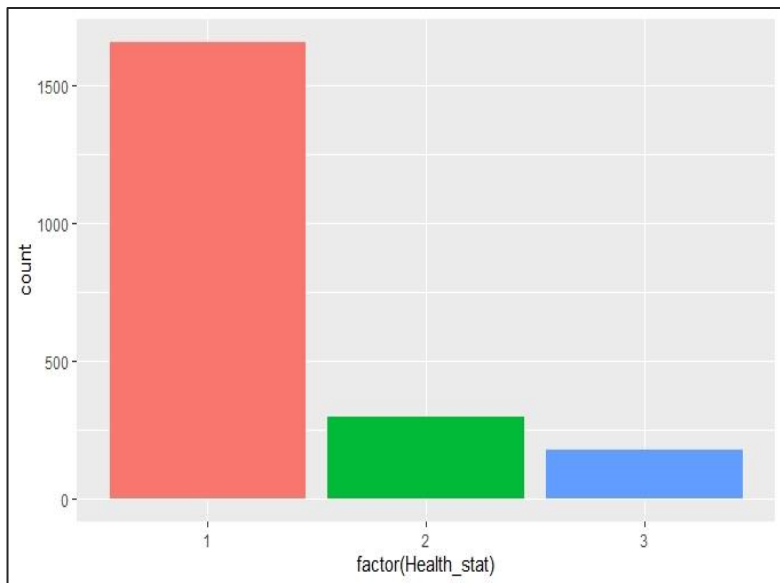


Data Structure

	LB	ACC	FM	UC	LD	SD	PD	LTV	ASV	MSV	MLV
min	106	0	0	0	0	0	0	12	0.2	0	0
max	160	0.019	0.481	0.015	0.015	0.001	0.005	87	7	91	50.7
range	54	0.019	0.481	0.015	0.015	0.001	0.005	75	6.8	91	50.7
median	133	0.002	0	0.004	0	0	0	49	1.2	0	7.4
mean	133.30	0.003178	0.009481	0.004366	0.001889	3.29E-06	0.000159	46.99012	1.332785	9.8466	8.187629
var	96.84222	1.49E-05	0.002178	8.68E-06	8.76E-06	3.28E-09	3.48E-07	295.592	0.78011	338.4	31.6771
std.dev	9.840844	0.003866	0.046666	0.002946	0.00296	5.73E-05	0.00059	17.1928	0.88324	18.39	5.62824

Data Structure

- Of the total counts 1655 (78%) fetuses were normal trace, 295 (14%) were suspected trace and 176 (8%) were pathological trace.
- Data suggests that there is a significant positive correlation between long term variability(+0.4), average short-term variability (+0.57) and fetal health status whereas accelerations have a significant negative correlation (-0.3) to fetal health status.

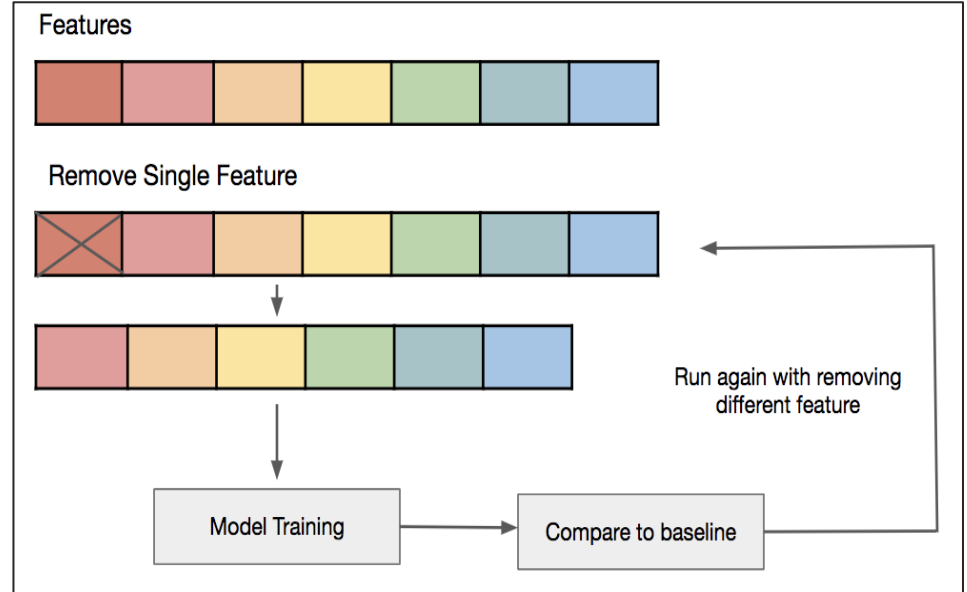


Algorithms used and model results

Feature Selection- Boruta^[8]

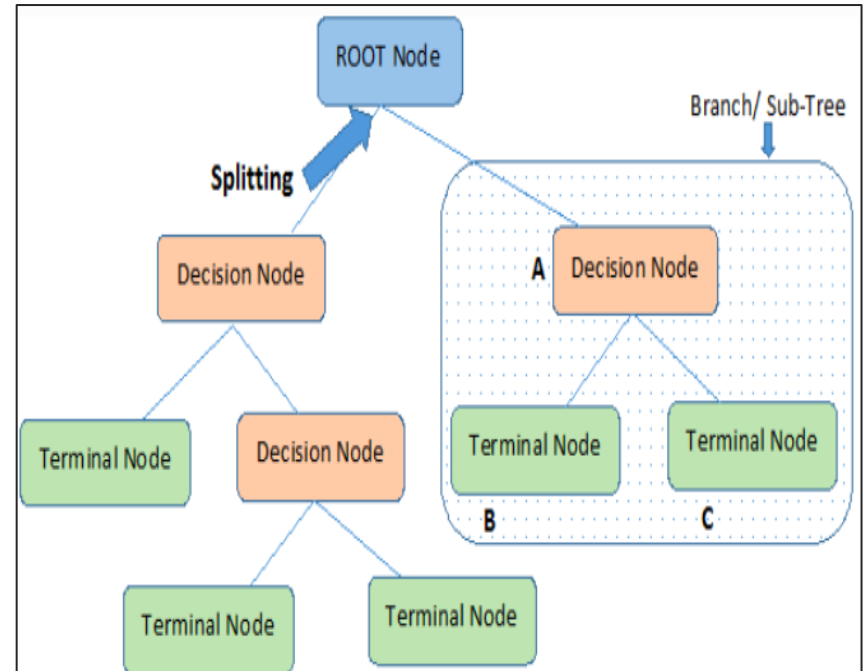
Boruta is a feature selection algorithm. Precisely, it works as a wrapper algorithm around Random Forest.

1. Firstly, it adds randomness to the given data set by creating shuffled copies of all features (which are called shadow features).
2. Then, it trains a random forest classifier on the extended data set and applies a feature importance measure (the default is Mean Decrease Accuracy) to evaluate the importance of each feature where higher means more important.
3. At every iteration, it checks whether a real feature has a higher importance than the best of its shadow features (i.e., whether the feature has a higher Z score than the maximum Z score of its shadow features) and constantly removes features which are deemed highly unimportant.
4. Finally, the algorithm stops either when all features gets confirmed or rejected or it reaches a specified limit of random forest runs.



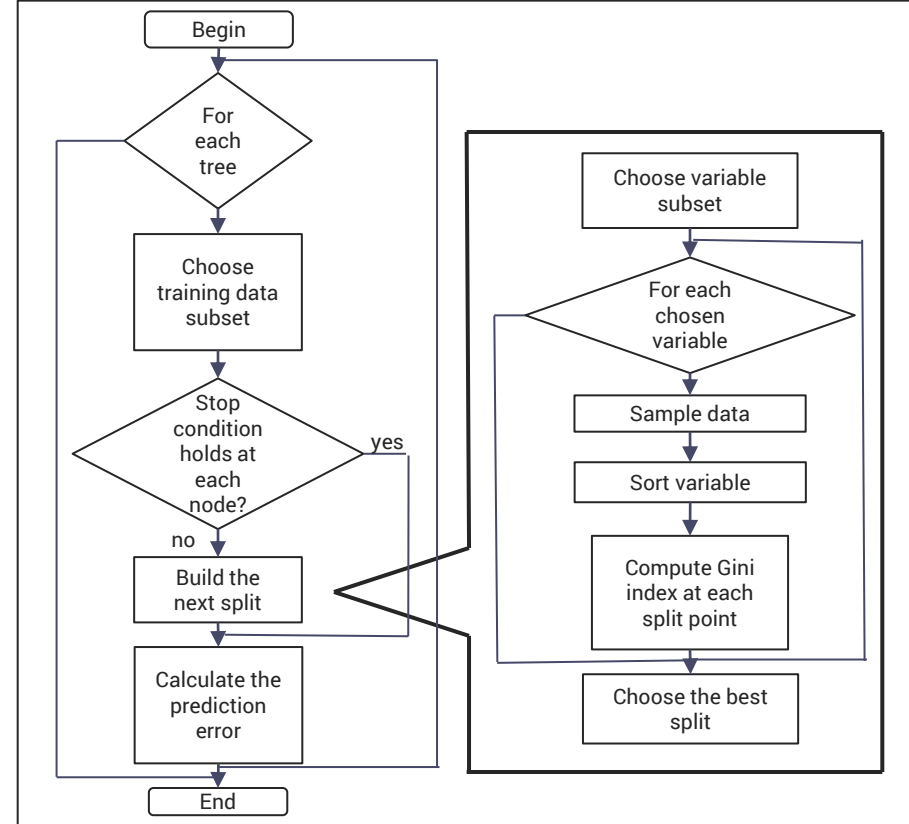
Decision Tree^[9]

1. A Decision tree is a flow chart like tree structure, where each internal node denotes a test on an attribute, each branch represents an outcome of the test, and each leaf node holds a class label.
2. A tree is constructed by splitting the source set into subsets on an attribute value test. This process is repeated in a recursive manner called as recursive partitioning. The recursion is completed when the subset at a node all has the same value of the target variable, or when splitting no longer adds value to the predictions.
3. Decision trees classify instances by sorting them down the tree from the root to some leaf node, which provides the classification of the instance. An instance is classified by starting at the root node of the tree, testing the attribute specified by this node, then moving down the tree branch corresponding to the value of the attribute.
4. This process is then repeated for the subtree rooted at the new node. The decision tree in the figure classifies a particular morning according to whether it is suitable for playing tennis and returning the classification associated with the particular leaf.(in this case Yes or No).



Random Forest^[10]

1. A random seed is chosen which pulls out a random collection of samples from the training set while maintaining the class distribution.
2. With selected data set, a random set of attributes from the original dataset is chosen based on defined values. All the input variables are not considered because of enormous computation size and chances of overfitting.
3. If there are “M” attributes in the whole data set, only “R” attributes are chosen at random for each tree where “R” < “M”.
4. The attributes from this set creates the best possible split using the gini index to develop a decision tree model. The process repeats for each of the branches until the termination condition stating that the leaves are the nodes that are too small to split.
5. Multiple trees are produced in a similar manner with different split and the outcomes are then aggregated to determine the target variable in a Random Forest.



Random Forest

Assuming there are n samples in the data with d features and classes $C1$ and $C2$ It can be represented by

$$D = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\}$$

Assuming there are S samples at the current node that need to be partitioned

$$P(S_j) = |S_j| / |S|$$

$$P(C_j/S) = |S_j \cap C_j| / |S_j|$$

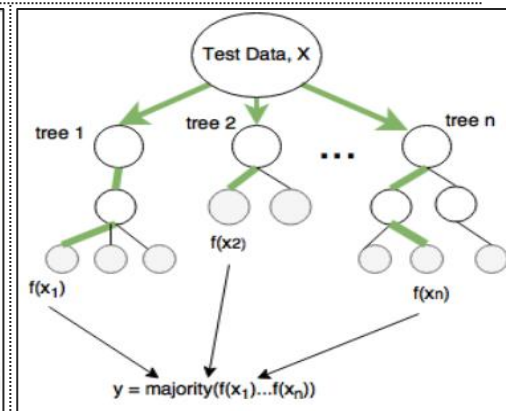
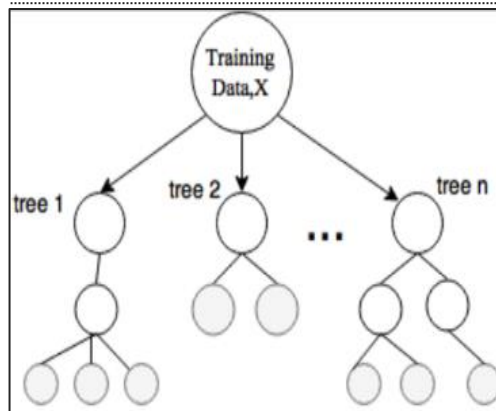
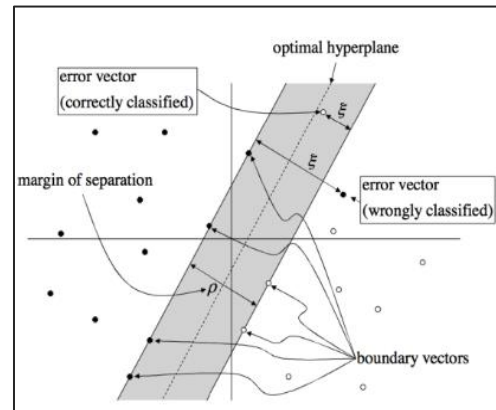
$$g(S) = \sum P(C_j/S)(1 - P(C_j/S))$$

Gini Index is given as,

$$G = P(S1) g(S1) + P(S2) g(S2)$$

Where $S1$ and $S2$ are the splits at each subsequent child node.

In a classification problem such as this the final output from majority of the trees is taken as the predicted value y . To prevent overfitting in the model splitting of trees should only be done till the point where the next split does not add to the overall accuracy in prediction.



Advantages of Random Forest

The model can be used for both classification and regression problems.

1. It can produce a highly accurate classifier and learning is fast.
2. It runs efficiently on a large data set.
3. It can handle thousands of input variables without variable deletion.
4. It computes proximities between pairs of cases that can be used in clustering, locating outliers or by scaling giving an overall view of the data.
5. It offers an exceptional means for detecting variable interactions.
6. Helps in identifying features based on their importance.
7. User friendly: the algorithm has only 2 free parameters
 - Trees (ntree) : Number trees used to develop the forest
 - Mtry: Variables randomly sampled as candidates at each split.

Training (input) model characteristics

	Predicted class 1 (Normal)	Predicted class 2 (Suspected)	Predicted class 3 (Pathological)
Actual class1	1160	2	0
Actual class 2	0	211	0
Actual class 3	0	0	122

Training statistics

- Training sample size: 1495
- Accuracy: 0.9987
- Accuracy at 95% Confidence Interval: (0.9952, 0.9998)
- No information rate: 0.7759
- p- value: <2.2e-16
- Kappa: 0.9964
- McNamara's Test P-Value: N/A

	(Normal)	(Suspected)	(Pathological)
Sensitivity	1.000	0.9906	1.000
Specificity	0.994	1.000	1.000
Positive prediction value	0.9983	1.000	1.000
Negative prediction value	1.000	0.9984	1.000
Prevalence	0.7759	0.1425	0.0816
Detection rate	0.7759	0.1411	0.0816
Detection prevalence	0.7773	0.1411	0.0816
Balanced accuracy	0.9970	0.9953	1.000

Training (input) model characteristics

1. The random forest consists of 400 trees to predict the fetal health outcome. Increase in the number of trees beyond this did not yield any significant increase in the predictive accuracy.
2. Out of bag error estimate of the model with 400 trees was 5.95% whereas for a model with 500 trees OOB was 6.62%

At 400 trees

```
randomForest(formula = Health_stat ~ ., data = train, ntree =  
300, mtry = 8, importance = TRUE, proximity = TRUE)
```

Type of random forest: classification

Number of trees: 400

No. of variables tried at each split: 8

OOB estimate of error rate: 5.95%

At 500 trees

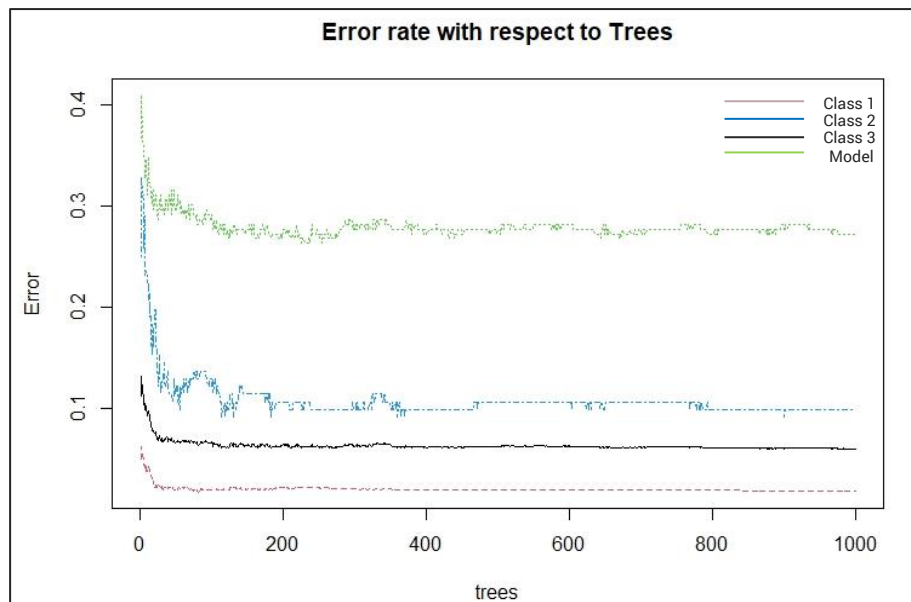
```
randomForest(formula = Health_stat ~ ., data = train, ntree =  
500, mtry = 8, importance = TRUE, proximity = TRUE)
```

Type of random forest: classification

Number of trees: 500

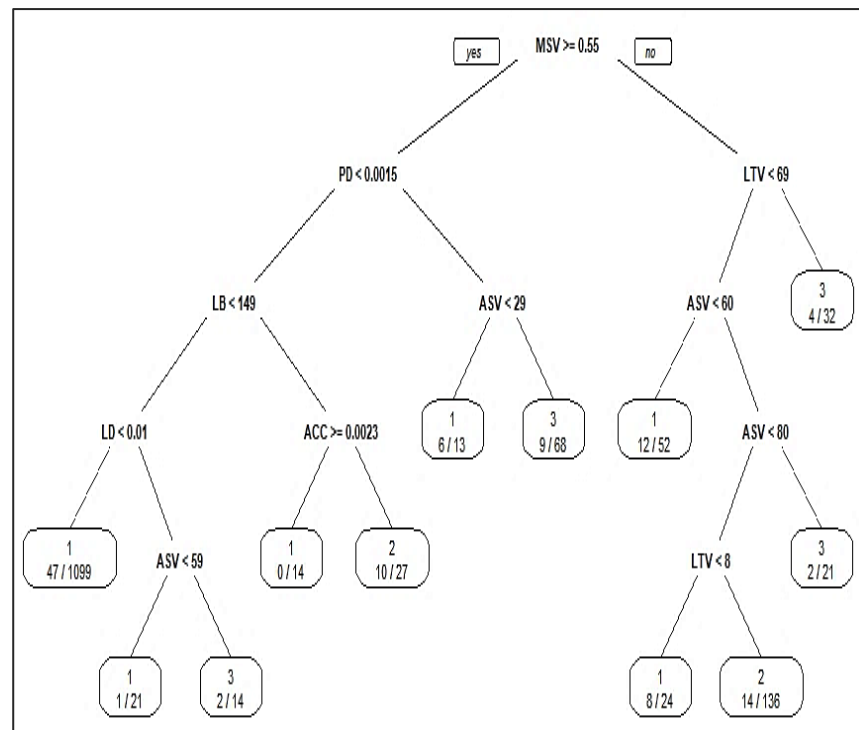
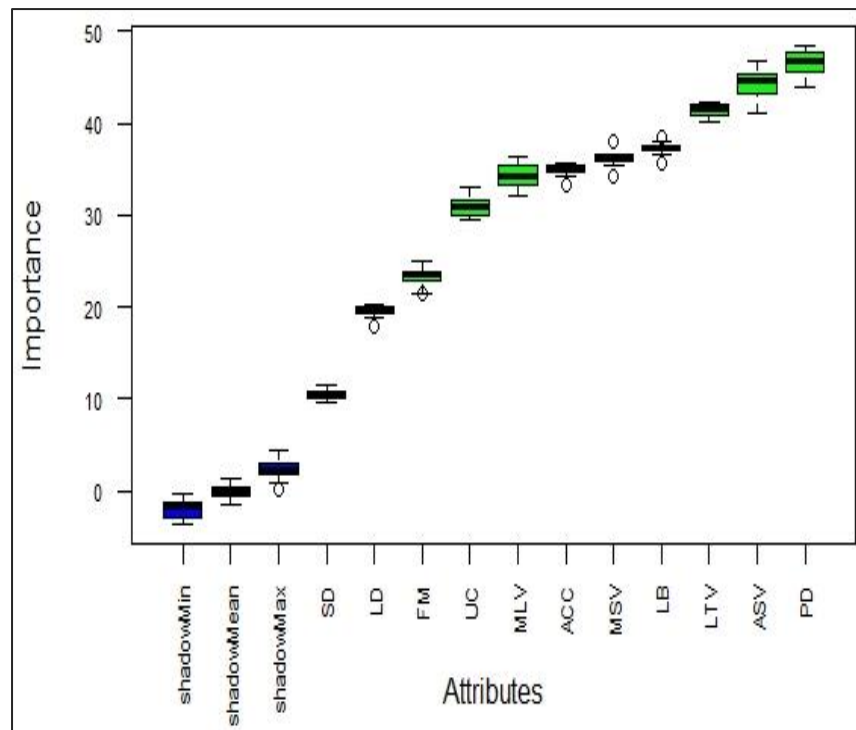
No. of variables tried at each split: 8

OOB estimate of error rate: 6.62%



Result- Boruta and Decision Tree

Boruta algorithm found all the 11 variables were important for the performance of the prediction model.



Result - Test (output) characteristics

	Predicted class 1 (Normal)	Predicted class 2 (Suspected)	Predicted class 3 (Pathological)
Actual class1	148	12	2
Actual class 2	16	61	5
Actual class 3	6	1	47

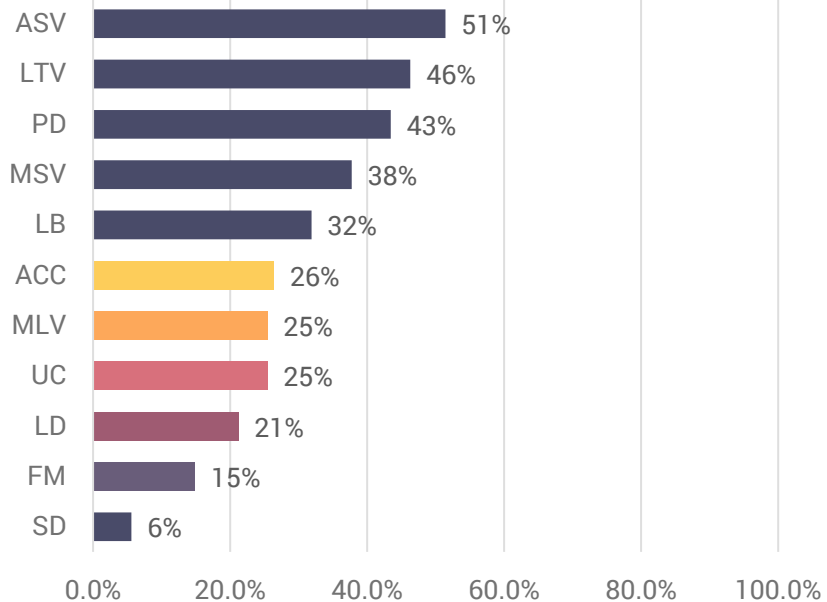
Test (Output) statistics:

- Accuracy: 0.9471
- Accuracy at 95% Confidence Interval: (0.9261, 0.9635)
- P-Value < 2.2e-16
- Kappa: 0.8453
- McNamara's Test P-Value: 0.004264

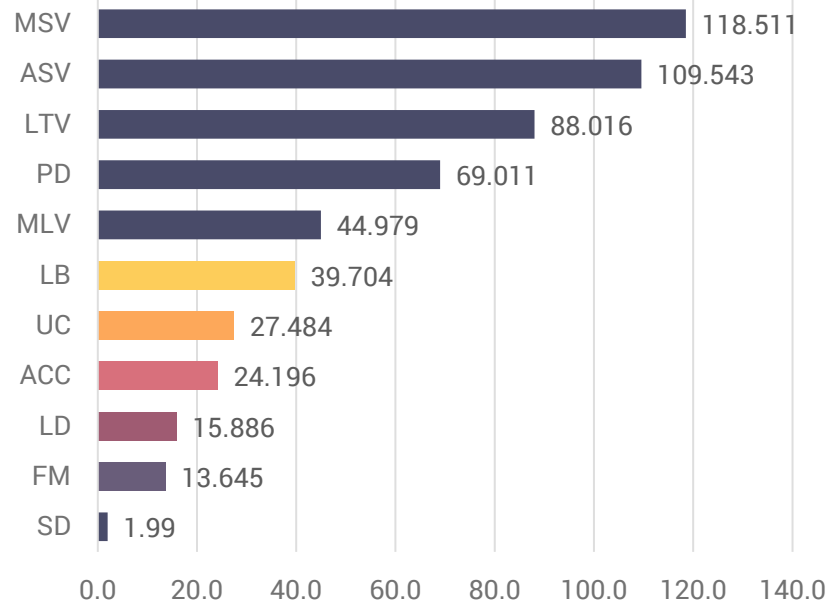
	(Normal)	(Suspected)	(Pathological)
Sensitivity	0.9873	0.7558	0.9853
Specificity	0.8244	0.9846	0.9982
Positive prediction value	0.9532	0.8904	0.9756
Negative prediction value	0.9474	0.9605	0.9913
Prevalence	0.7835	0.1421	0.0743
Detection rate	0.7736	0.1074	0.0661
Detection prevalence	0.8116	0.1207	0.0677
Balanced accuracy	0.9059	0.8702	0.9435

Result - Test (output) characteristics

Top 10 variables that effect accuracy



Top 10 variables by importance in the model



Research Results

1. Result #1

All features recorded during the non-stress test were important to develop the predictive model.

2. Result #2

The overall accuracy of the model stood at 94.7%, sensitivity and specificity for pathological trace is 0.98 and 0.99, respectively.



Comparison with similar studies

The model created in this study shows better predictive accuracy (94%) and sensitivity (0.98) for categorizing Pathological fetuses.

Sr. No	Reference	Method	Result (Category 3)	
			Accuracy	Sensitivity
1	Sundar et al 2013	Neural network-based classifier	0.91	0.90
2	Menai Mohder et al 2013	Relief F-15	0.939	0.958
3	Jezewski NSKI et al.	LSVM classifier	0.90	0.92
4	Chen et al 2012	FG- K means	0.76	0.81
5	Zhou and Sun 2014	Active learning of Gaussian process	0.89	0.79
6	Cruz et al 2016	Meta- DES Ensemble Classifier	0.846	-
7	This paper	Random forest with feature selection	0.94	0.98

Discussion and conclusions

Discussion

Discussion Point #1

Non-stress test plays a vital role in identifying high risk fetus during pregnancy and if appropriately managed can prevent fetal hypoxia and fetal death. The interpretation relies heavily on obstetrician ability to analyze traces and may lead to subjectivity.

Discussion Point #3

The strength of the study is that it uses Machine Learning algorithms on the data collected from non-stress test and developed a model that is high prediction accuracy for the pathological trace.

Discussion Point #2

Interpretation of the NST results requires an expert, ML based models can find application in remote areas where trained healthcare professionals are scarce. The role of eHealth technology can enhance the healthcare utilization in low- and middle-income areas improving the quality of antenatal care.

Discussion Point #4

More in-depth study and analysis can be put into accurately identifying the suspicious traces as the current model has a low predictive power for them.

Conclusions

How can Machine Learning based clinical decision support system be used to initiate an early response by identifying high risk pregnancies and prevent death due to fetal hypoxia?

The model developed in this study can predict the pathological fetal state with accuracy of 92-96%. Enabling the healthcare provider to ensure timely management.

The model developed in this study has a high accuracy(94%), sensitivity(0.98) and specificity(0.99) to identify pathological state than the once currently present.

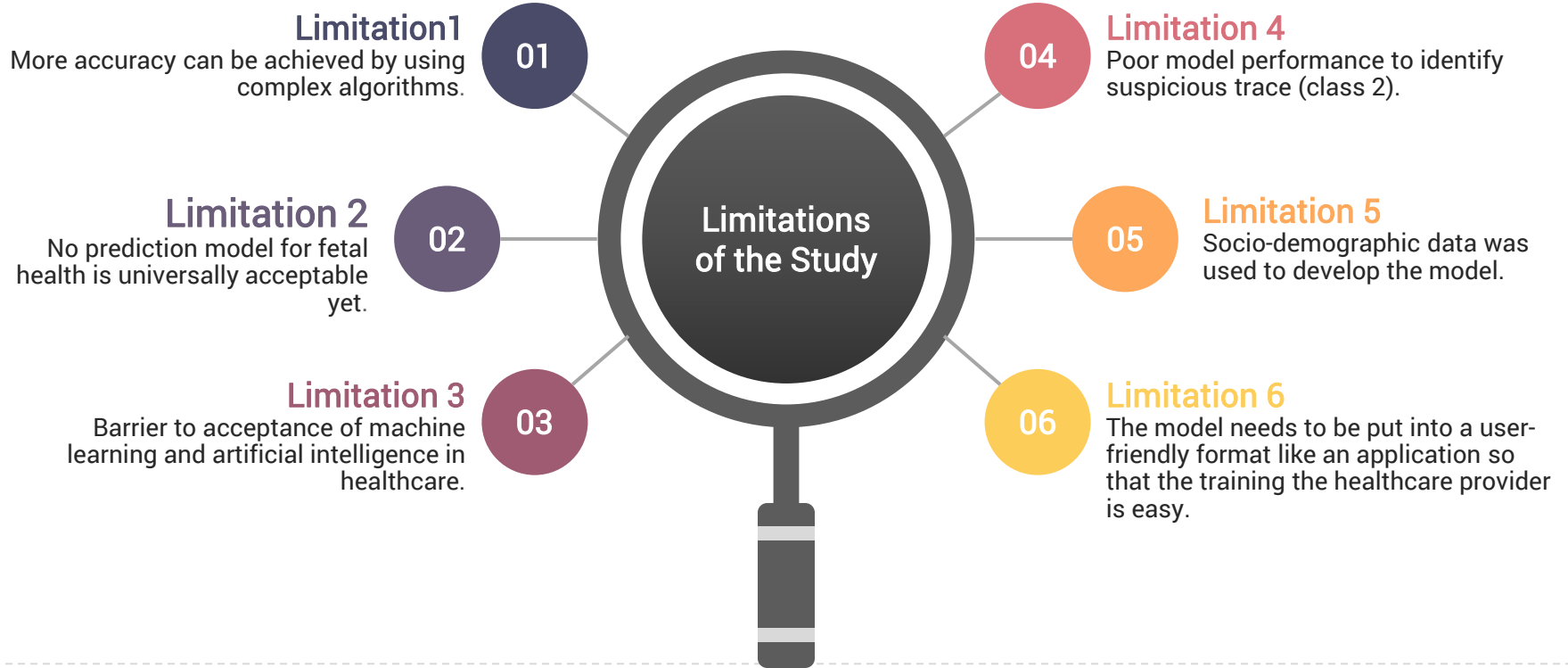
Some limitations:

1. The model needs to be put into a user-friendly interface.
2. Accuracy to identify suspicious trace is lower than the other classes

Further research is needed to incorporate socio-demographic features into this model.

Limitations of the study

Limitations of the Study



References

1 <https://data.unicef.org/country/ind>

2 <https://data.unicef.org/topic/child-survival/child-survival-sdgs/#:~:text=The%20proposed%20SDG%20target%20for,deaths%20per%201%2C000%20live%20births..>

3 <https://bmcpregnancychildbirth.biomedcentral.com/articles/10.1186/s12884-015-0472-9>

4 <https://bmcpregnancychildbirth.biomedcentral.com/articles/10.1186/s12884-015-0472-9>

5 <https://www.webmd.com/baby/nonstress-test-nst#1>

6 <https://www.mayoclinic.org/tests-procedures/nonstress-test/about/pac-20384577>

7 <https://www.ncbi.nlm.nih.gov/books/NBK537123/>

8 <https://towardsdatascience.com/boruta-explained-the-way-i-wish-someone-explained-it-to-me-4489d70e154a>

9 <https://www.geeksforgeeks.org/decision-tree>

10 <https://towardsdatascience.com/understanding-random-forest-58381e0602d2>

THANK YOU !