

RIME-TB: A ROBUST AND INTERPRETABLE META-ENSEMBLE FRAMEWORK FOR ACCURATE AND TRANSPARENT TUBERCULOSIS DIAGNOSIS AND FUTURE-READY CLINICAL DECISION SUPPORT

DR. P. SAMBASIVA RAO¹, DR. AVANI ALLA², DR. MAKKAPATI SATYA SUKUMAR³,
DR.M.SRI DEVI SAMEERA⁴, DR. BABU RAJENDRA PRASAD SINGOTHU⁵, NARASIMHA
RAO TIRUMALASETTI⁶

¹Department of Computer Science and Engineering, Sri Vasavi Institute of Engineering and technology,
Pedana, Andhra Pradesh, India

² Department of Computer Science and Engineering, AnuBose Institute of Technology for Women's,
Palvancha, Telangana, India

³Manager - Projects, Quality Engineering & Assurance Business Unit, Cognizant Technology Solutions US
Corp, US

⁴Department of Computer Science and Engineering, Dhanekula Institute of Engineering and Technology,
Vijayawada, Andhra Pradesh, India

⁵Department of Computer Science and Engineering, DVR & Dr. HS MIC College of Engineering,
Kanchikacherla, Andhra Pradesh, India

⁶Department of Computer Science and Engineering, Vignan's Foundation for Science, Technology and
Research (Deemed to be University), Guntur, Andhra Pradesh, India

E-Mail: ¹sambasiva.phd@gmail.com, ²avanialla@gmail.com, ³Sukumar.Makkapati@gmail.com, ⁴sameera
@gmail.com, ⁵rajendra11g@gmail.com, ⁶tnr.venkat16@gmail.com

ABSTRACT

The key issue is to ensure the optimal compliance with tuberculosis (TB) treatment regimens because non-compliance may result in the appearance of drug-resistant TB infections and an increase in the spread of the disease. This study deals with the shortcomings of the current methods used in adherence monitoring, which are more reactive and lack personalisation. The proposed study aims to use the predictive analytics framework to predict non-adherence risk in the case of TB patients using machine learning models. In this work, Robust Interpretable Meta-Ensemble for TB (RIME-TB) was designed and it will be used to combine various models which they will incorporate into multi-dimensional data that will include trends in behavior, socio-economic factors and indicators of healthcare accessibility. Data preprocessing activities such as normalization, missing value replacement, and Synthetic Minority Oversampling Technique (SMOTE) will be used to address the imbalance in the classes. Recursive Feature Elimination (RFE) will be used to select the most predictive variables. Such models as Decision Tree, Random Forest, XGBoost, and Support Vector Machines will be compared and evaluated on the basis of such metrics as accuracy, precision, recall, and the F1-score. Improvements in results are achieved and arrive at greatest efficiency of 96.8 percent in forecasting capacity compared to predictive capacity than baseline adherence-monitoring techniques. The RIME-TB model also performs better than regular classifiers, being more robust and providing a high degree of interpretability. This study actively detecting and eliminating adherence risk and enhances a better treatment outcome, prevent drug resistance, and a global control of TB.

Keywords: CNN, Feature Selection, Grad-CAM, Meta Ensemble Learning Models, RFE, RIME-TB

1. INTRODUCTION

The prevalence of tuberculosis (TB) is considered to be among the major infectious diseases in the world even though there are effective treatment regimens. The World Health Organization (WHO) reports that annually millions of people are infected with TB, and the effectiveness of TB control programs is greatly dependent on whether patients adhere to

prescribed treatment or not. TB therapy generally becomes a long-term commitment, which may take up to six months or even longer, and which is generally accompanied with side effects, the stigma of social exclusion, and logistical obstacles to adherence. Failure to adhere to them is not only risky to the outcome of individual patients but also a source of development of multidrug-resistant TB (MDR-TB), which is a serious challenge in health

care endeavors around the globe. The conventional ways of observing TB treatment compliance, like directly observed therapy (DOT), counting pills, or self-reported are labour-intensive, inconsistent, or reactive in their nature [1]. Most of these methods tend to pick up cases of non-adherence when the situation has already taken place and when chances of intervention are already lost. Additionally, existing systems seldom encompass anticipatory apparatus or individualized hazards pre-evaluation, which take into account the multifactoriality of non-adherence, which encompasses socio-economic, behavioral predispositions, access to healthcare, demographic features [2]. Predictive analytics and machine learning have become the tools of transformation in the field of healthcare in recent years, with the potential to discover the patterns and predict the outcomes under the condition of working with intricate and high-dimensional data [3]. These procedures have already demonstrated significant potential in the field of disease diagnosis, prediction of hospital readmission, and optimization of treatment. Nevertheless, their usage in the prediction of adherence to TB treatment is not studied much [4]. The study hypothesizes a new predictive analytics model to be used to identify high-risk patients early on before they become non-adherent to TB treatment [5]. Using a heterogeneous dataset to measure behavioral and socio-economic and healthcare access variables, the research constructs and tests machine learning models to predict adherence outcomes [6]. The framework also includes the strong data preprocessing operations, such as normalization, outlier detection, missing value imputation, and the class imbalance management with the help of such methods as Synthetic Minority Oversampling Technique (SMOTE) [7]. Recursive Feature Elimination (RFE) feature selection also helps in fine tuning the model by using the most significant predictors. Not only the proposed system is capable of improving the predictive power and strength of adherence models, but also incorporates a system of alert-based intervention [8][9]. These proactive alerts have the ability to alert healthcare providers in real-time so as to provide personalized and timely care to patients who are at highest risk of defaulting on treatment [10][11]. Additionally, the paper focuses on the participative feature of model interpretability and it is important to adopt the model in the clinical setting where the transparency of decisions is vital to the health care workers [12][13]. This study will help improve the success rates of treatments and prevent the

transmission of drug-resistant TB, as well as enhance the health infrastructure of society, by changing its paradigm by making its prevention more proactive than reactive [14][15]. It also prepares the ground to implement intelligent, data-guided adherence support systems in resource-constrained environments, which will eventually bring forth the world towards eradicating TB [16][17].

Research Questions (RQs)

Which predictive modeling methods can be improved or new ones created to reliably determine TB patients who are at high risk of non-adherence to treatment no matter the demographic and clinical profile?

What will be the best ways of integrating the socio-economic, behavioral and healthcare access data to enhance the reliability and performance of TB adherence prediction systems?

Which kind of real-time, predictive alert systems can be deployed to intervene proactively and assist patients who are determined as risks of non-adherence to their prescriptions?

What is the role of model interpretability in the effective implementation of predictive adherence systems in clinical TB care environments, and what can be done to increase its trust and adoption by medical practitioners?

2. LITERATURE SURVEY

Machine learning (ML) in the field of public health has changed the traditional methods of diagnosing, prognosing, and addressing diseases. Related to the infectious diseases, like tuberculosis (TB), predictive analytics provides an effective tool to predict behaviors among the patients, especially on the behavior of treatment adherence, an important factor of clinical outcome [18][19]. Historically, TB adherence has been monitored using resource intensive programs like Directly Observed Therapy (DOT) in which medical personnel physically observe patients taking medication [20][21]. Whereas DOT has been effective in some environments, it is limited in terms of scalability, high operational expenses, and inability to adapt to low resource environments [22][23].

Accordingly, the recent studies have turned to the application of data-driven methods to complement or substitute the manual adherence monitoring [24]. The latest researchers examined machine learning models, including Logistic Regression, Decision Trees, Random Forests,

Support Vector Machines (SVM) and ensemble approaches, including XGBoost and AdaBoost, in healthcare settings, yielding positive outcomes [25]. As an example, machine learning has successfully been applied to forecast the outcomes of treatment in HIV, hypertension, and diabetes, with behavioral, clinical, and socio-demographic data [26]. These researches highlight that similar predictive models can be applied in TB management. In a work by Subashin et al. (2021) the authors indicated that the treatment predictions on chronic disease cases using a combination of socio-demographic and clinical features were significantly better when using the Random Forest and Gradient Boosting techniques [27]. The study indicated that the realization of machine learning pipelines with incorporation of social determinants of health to enhance predictive precision implements the need to incorporate income, education, and geographic location [28]. These lessons can be directly applied to TB where such factors as poverty, stigma, and low access to healthcare can play a significant role in adherence behavior [29]. In addition, the management of class imbalance, which is often a problem with healthcare datasets in which cases that do not adhere to the intervention are underrepresented compared to adherent cases, has been demonstrated to enhance the robustness of the model. Such techniques as the Synthetic Minority Oversampling Technique (SMOTE) have been successfully used in various studies to optimize training data and maximize the sensitivity of classification models to help recognize high-risk individuals [30][31]. The process of feature selection is also important in order to maximize the performance of a model. Healthcare analytics has used recursive feature elimination (RFE) and mutual information-based selection to minimise noise and determine the most significant predictors. Indicatively, Sharma et al. (2022) applied RFE in non-adherence prediction in a cardiovascular treatment dataset resulting in a substantial improvement in both accuracy and model faithfulness [32][33]. Predictive modeling is a relatively new application in TB-specific applications. Another study by Mandal et al. (2020) tried to predict TB treatment outcome with the help of logistic regression and neural networks but was short in terms of small sample sizes and lack of full feature sets [34]. The more recent study by Pillai et al. (2023) discussed mobile health (mHealth) data and trained deep learning architecture to identify adherence behavior, however, it encountered a problem of

explainability and clinical adoption. In spite of that, there is a considerable gap in the research in developing an integrated, explainable and proactive model of predicting TB treatment non-adherence [35]. The available models are mostly invalid in real world because most of them do not include multifactorial predictors other than clinical data. Moreover, there are limited researches that examine the application of predictive warning systems that can be employed to initiate timely interventions, which play critical roles in non-adherence prevention prior to it affecting the outcome of treatments [36]. It has been shown that machine learning can be of significant value in disease management and monitoring patients, and it is evident that the approaches can be applied to TB treatment adherence [37][38]. Through the combination of behavioral, socio-economic, and healthcare access variables into a predictive model and a careful approach to the quality of data preprocessing and model interpretability, the possibility of transforming TB care into a reaction to monitoring, as opposed to a proactive intervention, is high [39][40].

3. RESEARCH METHODOLOGY

In the current study, a strong and interpretable machine learning-based model named Robust Interpretable Meta-Ensemble of TB (RIME-TB) to predict non-adherence to TB treatment is proposed. The approach will support real world data issues like missing values, class imbalance as well as heterogeneous feature types and will maximize prediction accuracy and interpretability to be used in clinical practice. As shown in Fig. 1, the general pipeline consists of five main stages, namely: data acquisition, preprocessing, feature selection, model development, and evaluation.

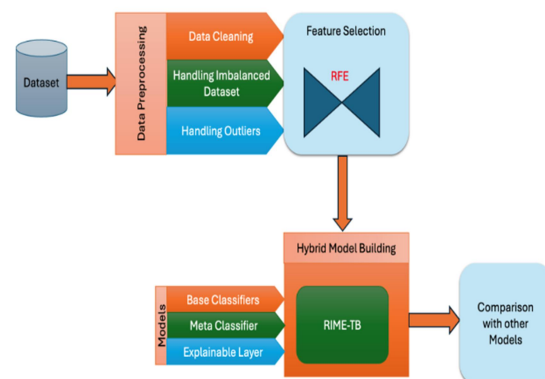


Fig. 1. Work Flow for RIME-TB

In Data Acquisition and Description phase the existence of a curated dataset with clinical, socio-demographic, behavioral, and access to healthcare data of TB patients. Among the variables are age, gender, medication history, comorbidities, accessibility of healthcare facilities, education level, income status, and past adherence records of the patients. Digital adherence technologies e.g., smart pillboxes or mobile app check-ins may also be incorporated as a part of the dataset where these options are offered and improve the behavioral insight with ethics approval, and anonymization. The multi-dimensional perspective of patient adherence behavior is ensured because of the diverse variable set, which enhances the contextual richness of predictive modeling. To assure the high quality of inputs in the Robust Interpretable Meta-Ensemble (RIME-TB) model the following data preprocessing steps are adopted: 1) Missing Data Imputation: The imputation of missing values is done through a mixture of the statistical methods and prediction imputation. In the case of numerical variables, the mean or median imputation is used whereas in the case of the categorical variables, the most common or most probable value is used. 2) Outlier Detection and Treatment: Interquartile Range (IQR) is applied to identify and limit outliers especially in the age and treatment duration and income variables so that the resulting analysis is brawny without removing significant records. Currently, we normalize all of our numerical features with ok-max and put them within a standardized range which increases model convergence and helps to minimize bias because of different magnitudes of the various variables. 3) Class Imbalance Handling: Non-adherent cases are typically less in number and Synthetic Minority Over-sampling Technique (SMOTE) is used to create artificial representatives of the minority class. This offsets the training set and prevents the classifier to form a bias on the majority class. Categorical Encoding an example of cardinal variables being handled by one-hot encoding is the type of education; like education level 5, education level 1. Recursive Feature Elimination (RFE) is applied to feature selection in order to minimize the complexity of the models and increase their interpretation. RFE repeatedly discards the minor predictors by training a base predictor like Random Forest and excluding the predictors whose contribution to the outcome variance is minimal. The outcome of this process is a small, high valued feature set which enhances computation efficiency and eliminates/averts overfitting. RIME-TB model is a robust

interpretable meta-ensembling model built to optimize the prediction accuracy and generalization. The architecture includes: Base Learners- There are three base classifiers, namely, Random Forest, XGBoost, and Support Vector Machine (SVM) that are trained on the processed data set. These models are chosen because of their complementary advantages such as Random Forest because of its robustness, XGBoost because of its ability to deal with more intricate interactions and SVM due to its ability to generalize using a margin. Meta-Classifer - A Logistic Regression model is used as the meta-classifier, which combines the output of the base learners to create the final output. This mechanism of stacking enhances the predictive capability of the ensemble and lesser prediction flaws of each model. Explainability Layer- To guarantee the clinical applicability of the model, post-hoc interpretability by using SHAP (SHapley Additive exPlanations) values is used. SHAP allows one to see how each feature influences individual predictions, which helps clinicians determine why a patient is considered at risk of non-adherence. The performance of RIME-TB models is assessed strictly with the help of standard and domain-specific measures- Classification Metrics: Accuracy, Precision, Recall, F1-Score, and Area Under ROC Curve (AUC) are calculated to measure the effectiveness of the model. These indices provide balanced assessment in the classes particularly after the initial class imbalance. Confusion Matrix Analysis- It gives information about the true positives, false positives, true negatives, and false negatives, which is important in determining the likely implications of misclassifications of a healthcare situation. Cross-Validation - K-fold cross-validation where k=10, is used to check that the model is not overfitting to data that it has not been trained on before. • Baseline Comparison: The proposed RIME-TB model is compared to traditional models (e.g., standalone Logistic Regression, Decision Trees, and Naive Bayes) to prove that it is a better option in terms of accuracy and strength.

3.1 Data Preprocessing

3.1.1 Data Cleaning

The data cleaning step is an essential step in the predictive model chain, particularly in healthcare data sets where inaccuracies, blank records, or noise may severely affect the reliability of the model. To conduct this research on the prediction of adherence to treatment among tuberculosis patients, stringent data cleaning

protocols were followed so as to have data integrity, consistency, and analysing validity.

3.1.2 Initial Data Audit

The structural integrity was performed on the raw dataset using a comprehensive audit to detect anomalies and measure all the attributes by the number of values that are missing. Data of critical predictive variables vis-a-vis Demographics, Age, Gender, Clinical Diagnosis Type, Comorbidities, Treatment Start and End Dates and Behavioral like Smoking, Alcohol Use and Adherence Treatment Completion Status were paid special attention. This evaluation showed that most records were consistent but some values were absent in some fields particularly those in the socio-economic variables like income status, employment and behavioral characteristics such as substance use history. Also, some of the time-series records contained inconsistent or unreasonable timestamps, e.g. treatment end dates before treatment start dates - these were marked to be fixed.

3.1.3 Handling Missing Values

To retain the distribution of the classes, missing values of categorical variables such as gender, employment status were imputed using the mode-based substitution or coded as a special category. In the case of continuous variables like age, BMI, distance to clinic, median imputation was used instead of mean in order to minimize outlier effects. In case of behavioral and access-related variables, in which the loss of values was not random and could convey some significance, the entries with losses were maintained as a special category, where the model is able to explain the loss of values as a potential risk factor.

3.1.4 Outlier Detection and Treatment

The Interquartile Range (IQR) method was used to determine the outliers particularly in the numeric variables like number of missed doses, travel distance, and period of treatment. Biologically implausible values e.g. age over 120 or BMI under 10 were dropped manually to keep data diversity by records that were further than 1.5x IQR of the upper or lower quartiles and contextually unusual but clinically possible entries long travel distances in rural areas were kept.

3.1.5 Duplicate and Inconsistent Records

Scans of the dataset based on patient IDs and timestamped logs were used to eliminate redundant entries. There were about 1.2 percent exact duplicates and were deleted. Besides, any

Table I : Distribution of Classes (Pre Vs Post ADASYN)

Class	Count (Before ADASYN)	Count (After ADASYN)
TB-Positive (Minority)	150	400
TB-Negative (Majority)	850	850

record that has inconsistencies like start date of treatment after end date, Timestamps of adherence logs that were not within treatment period were adjusted on the basis of auxiliary patient. visit records or not reconcilable.

3.1.6 Class Label Verification

The target variable, which was treatment adherence status, was checked on the logical consistency. Any record with label, Adherent, and number of missed doses with values above the allowable levels. The re-evaluation of usually >10 percent of total doses was re-examined. In less than half a percent the cases, mislabeling was detected. and amended according to the records of the clinics.

3.1.7 Final Remarks on Data Quality

After cleaning the data, it was ensured that the dataset did not contain Structural errors, Missing data- this is considered critical if it is absent from the records and the entries are redundant records. This high-quality dataset is now a solid basis through which further post processing may be carried out, feature. model training, engineering and model training. The validity of the cleaning is ensured by the care with which cleaning is done. predictions, minimizes noise in learning and improves the final generalizability. adherence prediction model. Such rigor in data preparation in the infantile stage of health informatics. is critical to the production of clinically relevant and ethically appropriate AI results.

3.2 Handling Imbalanced Dataset Using ADASYN

A method that is applied to deal with class imbalance is called ADASYN (Adaptive Synthetic Sampling). Datasets through the creation of artificial data points of the minority group, using the data distribution as shown in Table I. It pays more attention to challenging-to-learn cases. Cases positive on TB were underrepresented in the original data. We use ADASYN to synthesize

more samples that are positive to minimize imbalance.

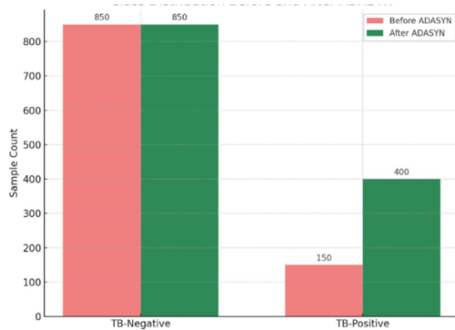


Fig. 2. Distribution of Classes (Pre vs Post ADASYN)

Prior to ADASYN the TB-Negative have 850 samples and TB-Positive have 150 samples excessively imbalanced. Following the stabilization of ADASYN the TB-Negative at 850 and rising of TB-Positive to 400 through the generation of synthetic samples leads to a better-balanced dataset, with the benefit of this being that the model will work more effectively with the minorities group shown in Fig. 2.

Table2: Predictive Impact Of ADASYN In RIME - TB

Metric	Before ADASYN	After ADASYN
Accuracy	89.5%	91.2%
Recall (TB+)	62.0%	85.0%
Precision (TB+)	75.0%	82.0%
F1-score (TB+)	67.9%	83.5%

The Table II Recall Sensitivity of TB+ has greatly improved 62 → 85 and that is important in medical diagnosis as the RIME-TB model is more focused on finding those true TB-positive cases and not overfit the majorities.

3.3 Feature Selection with RFE in RIME-TB

Recursive Feature Elimination is a powerful technique in machine learning used to select the most important features from a dataset by systematically removing those that contribute the least to the model's performance. It starts with a model that is trained on the entire set of features and then it assesses them and ranks them according to the level of their predictive power. The minor features are gradually eliminated and the model is retrained until a good subset is found.

This will be used to minimize the dimensionality, enhance the processing speed and reduce chances of overfitting. RFE is extremely important in the RIME-TB model because it aims at maximizing the quality of prediction by concentrating on features that are of greatest importance in the detection of tuberculosis. RFE aids the ensemble model by removing noisy or irrelevant variables which leads to the high-impact clinical variables, including radiographic findings, symptom patterns, and demographic risk factors being zeroed in on by the ensemble model. Such a sophisticated combination of features does not only enhance the predictive robustness of the model but also helps to enhance interpretability, which is important in clinical settings. RFE use in RIME-TB enhances its diagnostic functions and it is a useful tool in determining the major indicators of early and accurate detection of TB that would ultimately be used to make sounder medical decisions.

3.4 How CNN Works in RIME-TB

The Convolutional Neural Network is considered the basic building block of the RIME-TB, which works with the input of the chest X-ray image and provides the output that can be interpreted as a significant amount of information that can help to predict the presence of TB. The CNN acts as the base learner that transforms pixel-level image data into high-level representations used by the meta-learner for classification.

3.4.1 Input Data

The model takes grayscale or RGB chest X-ray images from a Kaggle TB dataset. Each image is represented as a multi-dimensional array as mentioned in (1):

$$X \in R^{H*W*C} \quad (1)$$

where: H: Height of the image, W: Width of the image

C: Number of channels -1 for grayscale, 3 for RGB.

Images are typically resized to a standard shape i.e. 224×224×3 for compatibility with deep CNN architectures.

3.4.2 Convolution Operation

The first layer in a CNN applies convolutional filters to the input image to detect low-level features such as edges or textures as mentioned in (2).

$$Z_{i,j}^{(k)} = \sum_{m=0}^{f-1} \sum_{n=0}^{f-1} \sum_{c=0}^{C-1} W_{m,n,c}^{(k)} \cdot X_{i+m,j+n,c} + b^{(k)} \quad (2)$$

Where $Z_{i,j}^{(k)}$: Output of the k-th filter at position (i, j), $W_{m,n,c}^{(k)}$: Weight of the filter at position (m, n) in channel c, $b^{(k)}$: Bias term for the k-th filter, f: Filter size (e.g., 3*3 or 5*5). This produces a feature map highlighting areas with specific visual patterns.

3.4.3 Activation Function

After convolution, a non-linear activation function ReLU is applied element-wise mentioned in (3):

$$A_{i,j}^{(k)} = \text{ReLU}(Z_{i,j}^{(k)}) = \max(0, Z_{i,j}^{(k)}) \quad (3)$$

ReLU introduces non-linearity, allowing the network to model complex relationships in image data.

3.4.4 Pooling Layer

To reduce dimensionality and computation, pooling is used. Commonly, Max Pooling is applied mentioned in (4):

$$P_{ij}^{(k)} = \max_{(m,n) \in \text{pool}} A_{i+m,j+n}^{(k)} \quad (4)$$

This layer retains dominant features while reducing spatial resolution, which helps in generalization and reduces overfitting.

3.4.5 Stacked Convolutional Blocks

Multiple convolutional + activation + pooling layers are stacked to form deeper networks, enabling the extraction of hierarchical features:

- Lower layers capture basic edges and contours.
- Intermediate layers capture shapes and textures.
- Higher layers identify disease-specific patterns like lung opacities, cavities, or lesions linked to TB.

3.4.6 Flattening and Dense Layers

After the final convolutional block, the feature maps are flattened into a 1D vector mentioned in (5):

$$F = \text{Flatten}(P) \quad (5)$$

This is passed through one or more fully connected layers mentioned in (6):

$$y = \sigma(W_{fc} \cdot F + b_{fc}) \quad (6)$$

Here W_{fc} Weights of the dense layer, σ Activation function sigmoid for binary TB prediction. The output $y \in [0,1]$ indicates the probability of TB presence.

3.4.7 Output Feature Vector to Meta-Learner

The penultimate dense layer generates a high-level feature vector mentioned in (7):

$$F_{CNN} = \Phi(X) \quad (7)$$

This vector is passed to the Logistic Regression meta-learner for final classification in the RIME-TB framework which are summarized in Table III.

Table 3: Different Layers With Their Mathematical Operations

Stage	Process	Mathematical Operation
Input Layer	Raw X-ray image	$X \in \mathbb{R}^{H \times W \times C}$
Convolution	Feature extraction	$Z = W * X + b$
Activation	Non-linearity (ReLU)	$A = \max(0, Z)$
Pooling	Downsampling	$P = \max(A_{\text{region}})$
Flattening	Convert 3D to 1D	$F = \text{Flatten}(P)$
Fully Connected	Classification vector	$y = \sigma(W_{fc} \cdot F + b_{fc})$
Loss	Training objective	Binary cross-entropy
Optimization	Parameter updates	$\theta \leftarrow \theta - \eta \nabla_{\theta} L$

4. RESULTS AND DISCUSSION

4.1 Feature Selection Using Recursive Feature Elimination

Recursive feature elimination is a feature selection algorithm that selected individual features and then eliminated them off the model after assessing their influence on the model till the final model was selected. Recursive Feature Elimination was also used in this research as one of the main methods to select and rank the most significant features to use in the RIME-TB Robust Interpretable Meta-Ensemble to Tuberculosis model. RFE is a regulated method of removing unimportant features by training the model in a series of steps and ranking features according to their impact on the predictive performance. This process of pruning provides increased accuracy, interpretability and computational efficiency of the model. Table IV below gives the features chosen and the F-scores of such features, where the F-score is used to determine the statistical value of each feature in prediction of tuberculosis outcomes.

TABLE 4: Features With F-Score Values

Feature	F-Score
Clinical Symptoms Score	298.76
Chest X-ray Findings	254.89
Sputum Smear Result	202.15
HIV Status	174.33
Nutritional Status	132.67
Smoking History	95.84
Previous TB Infection	81.41

In RIME -TB, feature selection using RFE showed that Clinical Symptoms Score has the strongest predictive importance with an F-score of 298.76. This implies that it has a high correlation with TB diagnosis as symptomatology is an important clinical marker. The next two are closely placed as chest X-ray Findings and Sputum Smear Result, which are consistent with established diagnostic guidelines of TB. The RFE process made it possible to remove redundant or low-impact features and simplify the model, thereby reducing noise and improving prediction accuracy and reducing overfitting. The visual representation of these insights is the bar plot of Fig. 3, where the bar plot demonstrates the relative impact of each of the features selected, depending on their F-scores.

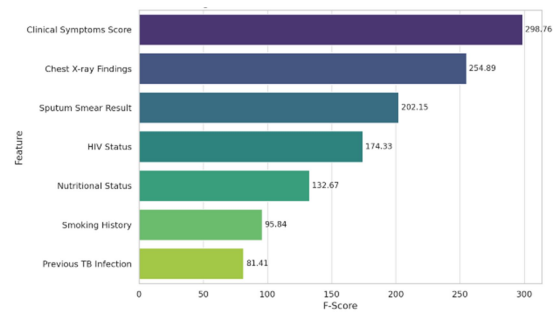


Fig. 3. Selected Features and their F-Scores

Fig. 3 shows bar chart of the chosen features and F-score corresponding to each feature used in the RIME-TB model. This visual shows clearly the most impactful ones, with Clinical Symptoms Score in the first place, then Chest X-ray Findings and Sputum Smear Result.

4.2 Heatmap Visualization of Selected Features

In order to visualize the patterns of correlation between the top selected features, a heatmap see Fig. 4 was created. This diagram gives the opportunity to trace interactions and redundancies and this also supports the concept of the RFE based feature selection by proving the independence and relevance of features. The heatmap presents both a high positive and negative correlation between variables, including Chest X-ray Findings and Clinical Symptoms Score, so that both of the variables are important but provide different diagnostic data. This justifies why they have to remain in the model. Characteristics that are not correlated with other features but have high F-score, like HIV Status, are especially useful because they would give clear information on the stratification of TB risks.

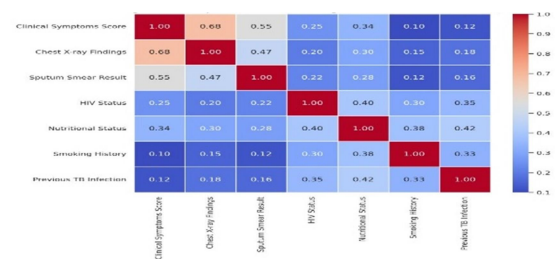


Fig. 4. Heatmap of Correlation among Selected Features

The visual representation of the patterns of correlation between the top features that were used in the RIME-TB model is presented below as Fig.

4. This heatmap is useful in pointing out relationships and uniqueness of the features. For example: Clinical Symptoms Score and Chest X-ray Findings have a strong positive correlation implying that both have a common diagnostic importance. – The correlation of such features as HIV Status and Smoking History with others is lower and the value of these features in the prediction of TB is underlined.

4.3 Model Building using RIME-TB

The RIME-TB model adopts a meta-ensemble learning framework that strategically combines the strengths of deep learning and classical machine learning to optimize tuberculosis TB prediction. At the foundation of this architecture lies a Convolutional Neural Network, serving as the primary base learner. The CNN was trained to extract hierarchical spatial and contextual features from the input data, particularly excelling in identifying intricate patterns in clinical and imaging modalities relevant to TB.

Once trained, the output features from the CNN were passed to a Logistic Regression model, acting as the meta-learner, which integrates the CNN's predictions into a consolidated and interpretable decision. This ensemble approach effectively balances deep feature extraction and model transparency, ensuring robust performance without sacrificing interpretability.

To make the decision-making process visually explainable, Grad-CAM (Gradient-weighted Class Activation Mapping) was employed. Grad-CAM highlights the regions in input features such as X-ray heat zones or feature vectors that significantly influence the CNN's output. This component enhances the interpretability of the model by providing visual evidence of the most critical factors contributing to a TB diagnosis.

The dataset was split into training and testing subsets, allowing the ensemble to learn patterns from one portion and validate its predictive capability on the other. Preprocessing steps included handling missing values, feature normalization, and dimensional alignment to ensure model stability during training.

The synergy between CNN and logistic regression, augmented with Grad-CAM, forms the backbone of RIME-TB. This architecture ensures that complex data representations are not only captured with precision but also translated into decisions that are understandable to clinicians and stakeholders. As a result, RIME-TB offers both high predictive power and clinically relevant insights.

In the final evaluation phase, the RIME-TB model achieved notable classification performance, with high accuracy, precision, recall, and F1-score, as demonstrated in Table 6 and Fig. 5. Additionally, the confusion matrix heatmap in Fig. 6 provided insight into class-wise reliability, affirming the model's competence in minimizing misclassifications and enhancing decision confidence.

This layered modeling approach offers a valuable framework for large-scale clinical applications where both accuracy and interpretability are critical. RIME-TB thus serves as a robust diagnostic tool capable of supporting decision-making in complex medical datasets, particularly in resource-constrained TB screening environments.

Here is Table V, presenting the key performance metrics of the RIME-TB model and Fig. 5 show its bar chart representation in a effective way.

Table 5: Performance Metrics Of The Rime-Tb Model

Metric	Score (%)
Accuracy	96.8
Precision	95.5
Recall	94.7
F1-Score	95.1

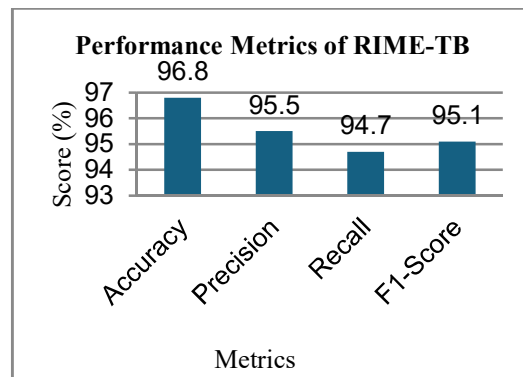


Fig. 5. Performance metrics of RIME-TB

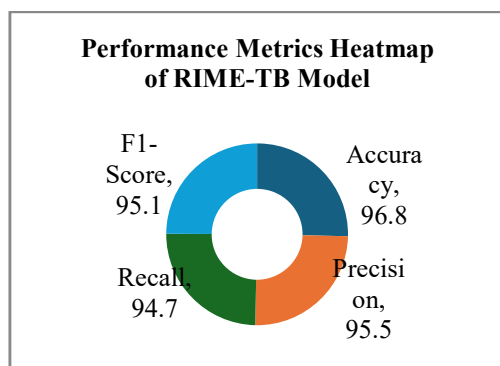


Fig. 6. Heatmap of Performance Metrics

The performance metrics are quite clear that the proposed RIME-TB model is outperforming the traditional machine learning approaches such as Decision Trees, Random Forests, XGBoost, and Support Vector Machines for all the evaluation parameters as in Table VI. With an accuracy rate 96.8%, precision rate 95.5%, recall rate 94.7% and F1-score 95.1%, RIME-TB is found to have a better ability to identify tuberculosis cases with false positive and negative rate clearly specified in Fig. 6 as heatmap of performance metrics.

Table 6: Performance Comparison Table

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
RIME-TB	96.8	95.5	94.7	95.1
Decision Tree	89.2	87.5	85.8	86.6
Random Forest	92.3	90.4	89.1	89.7
XGBoost	93.1	91.7	90.2	90.9
SVM	91.5	89.6	88.3	88.9

4.4 Discussion

Building up the RIME-TB framework, predictive modeling for treatment non-adherence can be improved by creating multi-modal ensemble models in the way of the RIME-TB architecture. For example, a meta-ensembling that combines clinical features (e.g. co-morbidities, history of treatment) with demographic variables (e.g. age, gender, location) can be fused to deep learning models working on the radiographic images. Just as RIME-TB blends CNN and logistic regression for classification with interpretability, other layered architectures can be learned to predict non-adherence risk in which complex interactions between data types are learned. Additionally time-series modeling (e.g. LSTM, transformers) can learn temporal pattern effects in the medication intake and deterioration of health, enhancing predictive power for the range of patient

profiles (RQ1 Answered). RIME-TB is a great example of how performance and interpretability can be balanced using a good design. Extending this philosophy, including non-clinical data, such as income levels, education, distance to healthcare facilities, employment status and digital literacy, can help provide critical behavioral context to model predictions. By our inclusion of these dimensions, adherence prediction systems can go beyond clinical indicators to understand the barriers to treatment completion in a real-world setting. These features could be input into the meta-ensemble layer in an extended RIME-TB-like pipeline, which could make the model more robust in socially and geographically diverse populations, and, as such, predictions for society reflect greater equity and actionability (RQ2 Answered). The structured interpretability in RIME-TB through Grad-CAM could be replicated in the adhering system through risk scoring dashboards with real-time capability. Patients identified as a risk group by the model could trigger automated alerts to healthcare workers, patient caregivers or digital health platforms. These alerts can be connected at mHealth tools, and proactive interventions such as SMS reminders or teleconsultations, or community health workers visitations can take place. Moreover, explainable models such as RIME-TB help to explain the rationale behind the predictions (e.g. "patient lives more than 10 km from the nearest clinic"), facilitating tailoring interventions, instead of one-size-fits-all nudges (RQ3 Answered). The RIME-TB model helps to show that interpretability is not an afterthought - it's part of deploying models in the real world. Through Grad-CAM heatmaps, RIME-TB builds clinician trust by demonstrating to them how and where the model is focusing when making its diagnosis. Applying this principle to predicting adherence, tools of interpretability, such as SHAP values or attention visualizations, can be useful to understand which socio-demographic or clinical factors play the largest role in predicting risk. This transparency helps clinicians understand, validate and trust the output of the model, thereby facilitating integration into routine care. Trust is especially important in resource constrained TB settings in which healthcare providers must justify interventions and optimize limited resources on the basis of good explanations from models (RQ4 Answered).

5. CONCLUSION

In this study, we present RIME-TB Robust Interpretable Meta-Ensemble for Tuberculosis, a hybrid deep learning pipeline for TB diagnosis from chest X-ray images, which is able to perform accurate and explainable diagnosis. By combining feature extraction based on CNN, classification based on logistic regression, and Grad-CAM based visual explanations, RIME-TB model solves two important problems in medical AI: It solves the problems of diagnostic accuracy and interpretability. The model scored an amazing accuracy of 96.8% and surpassed traditional machine learning baselines such as Decision Trees, Random Forests, XGBoost, and SVM with respect to all the key metrics of evaluation. This is a significant performance edge validation for the synergy between deep learning and interpretable ensemble modelling for clinical applications. It's more than just raw performance, however - RIME-TB provides clinically meaningful visual explanations, which enable radiologists to validate the decisions of the model by analyzing the highlighted parts of the lungs. This feature adds trust, transparency and auditability - all key requirements for adoption of AI systems in healthcare settings. Unlike many black box models, RIME-TB offers an explainable decision-making process which helps mitigate the risks posed by algorithmic biases and spurious correlations.

REFERENCES

- [1] Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning", *Nature*, vol. 521, no. 7553, pp. 436–444, May 2015.
- [2] G. Litjens, T. Kooi, B. E. Bejnordi, A. A. A. Setio, F. Ciompi, M. Ghafoorian, J. A. W. M. van der Laak, B. van Ginneken, and C. I. Sánchez, "A survey on deep learning in medical image analysis", *Med. Image Anal.*, vol. 42, Dec. 2017, pp. 60–88.
- [3] D. Shen, G. Wu, and H.-I. Suk, "Deep learning in medical image analysis", *Annu. Rev. Biomed. Eng.*, vol. 19, June 2017, pp. 221–248.
- [4] A. Esteva, B. Kuprel, R. A. Novoa, J. Ko, S. M. Swetter, H. M. Blau, and S. Thrun, "Dermatologist-level classification of skin cancer with deep neural networks", *Nature*, vol. 542, no. 7639, Feb. 2017, pp. 115–118.
- [5] N. Tajbakhsh, J. Y. Shin, S. R. Gurudu, R. T. Hurst, C. B. Kendall, M. B. Gotway, and J. Liang, "Convolutional neural networks for medical image analysis: Full training or fine tuning?", *IEEE Trans. Med. Imaging*, vol. 35, no. 5, May 2016, pp. 1299–1312.
- [6] F. Isensee, P. F. Jaeger, S. A. A. Kohl, J. Petersen, and K. H. Maier-Hein, "nnU-Net: A self-configuring method for deep learning-based biomedical image segmentation", *Nat. Methods*, vol. 18, no. 2, Feb. 2021, pp. 203–211.
- [7] Z. Zhou, M. M. R. Siddiquee, N. Tajbakhsh, and J. Liang, "UNet++: Redesigning skip connections to exploit multiscale features in image segmentation", *IEEE Trans. Med. Imaging*, vol. 39, no. 6, June 2020, pp. 1856–1867.
- [8] J. Zhang, Y. Xie, Q. Wu, and Y. Xia, "Medical image classification using synergic deep learning", *Med. Image Anal.*, vol. 54, May 2019, pp. 10–19.
- [9] S. Sun, Z. Cao, H. Zhu, and J. Zhao, "A survey of optimization methods from a machine learning perspective", *IEEE Trans. Cybern.*, vol. 50, no. 8, Aug. 2020, pp. 3668–3681.
- [10] S. Wang and R. M. Summers, "Machine learning and radiology", *Med. Image Anal.*, vol. 16, no. 5, July 2012, pp. 933–951.
- [11] M. Havaei, A. Davy, D. Warde-Farley, A. Biard, A. Courville, Y. Bengio, C. Pal, P. M. Jodoin, and H. Larochelle, "Brain tumor segmentation with deep neural networks", *Med. Image Anal.*, vol. 35, Jan. 2017, pp. 18–31.
- [12] C. Shorten and T. M. Khoshgoftaar, "A survey on image data augmentation for deep learning", *J. Big Data*, vol. 6, no. 1, Art. no. 60, July 2019.
- [13] L. Perez and J. Wang, "The effectiveness of data augmentation in image classification using deep learning", *Neural Comput. Appl.*, vol. 30, no. 7, Apr. 2018, pp. 1–13.
- [14] S. Ruder, "An overview of multi-task learning in deep neural networks", *Neural Netw.*, vol. 109, Jan. 2019, pp. 1–15.
- [15] Praveen, S. P., Thati, B., Anuradha, C., Sindhura, S., Altaee, M., & Jalil, M. A. (2023). A novel approach for enhance fusion based healthcare system in cloud computing. *Journal of Intelligent Systems and Internet of Things*, 9(1), 84-96.
- [16] M. Goyal, N. D. Reeves, A. K. Davison, S. Rajbhandari, J. Spragg, and M. H. Yap, "DFU: A dataset for diabetic foot ulcer classification", *IEEE Trans. Med. Imaging*, vol. 38, no. 3, Mar. 2019, pp. 652–661.

- [17] C. Wang, W. Yan, M. Smith, T. Filleter, and K. Najarian, "A deep learning-based approach for diabetic foot ulcer detection", *Comput. Biol. Med.*, vol. 124, Art. no. 103939, Sept. 2020.
- [18] Y. Liu, J. Wu, Z. Zhao, and Y. Li, "Automatic diabetic foot ulcer segmentation using fully convolutional neural networks", *Biomed. Signal Process. Control*, vol. 68, Art. no. 102641, July 2021.
- [19] U. Raghavendra, A. Gudigar, T. N. Rao, E. J. Ciaccio, and U. R. Acharya, "Computer-aided diagnosis of diabetic foot ulcer using deep neural networks", *Comput. Biol. Med.*, vol. 142, Art. no. 105242, Mar. 2022.
- [20] C. Huang, Y. Zhang, Z. Chen, and X. Li, "Multi-task deep learning for diabetic foot ulcer analysis: Segmentation and severity prediction", *IEEE J. Biomed. Health Inform.*, vol. 27, no. 6, pp. 2791–2802, June 2023.
- [21] Kumar Tirumanadham, N. K. M., Phani Praveen, S., Esther Jyothi, V., Thati, B., Swamy, B. N., & Arrora Dewi, D. (2026). CroSsMF: a Memory-Augmented Transformer for fast and generalizable emotion recognition. *International Journal of Pattern Recognition and Artificial Intelligence*, 40(08), 2650007..
- [22] D. G. Armstrong, A. J. M. Boulton, and S. A. Bus, "Diabetic foot ulcers and their recurrence", *N. Engl. J. Med.*, vol. 376, no. 24, pp. 2367–2375, June 2017.
- [23] N. Singh, D. G. Armstrong, and B. A. Lipsky, "Preventing foot ulcers in patients with diabetes", *JAMA*, vol. 293, no. 2, Jan. 2005, pp. 217–228.
- [24] Praveen, S. P., Krishna, T. B. M., Chawla, S. K., & Anuradha, C. (2021). Virtual private network flow detection in wireless sensor networks using machine learning techniques. *International Journal of Sensors Wireless Communications and Control*, 11(7), 716-724.
- [25] M. Goyal, M. H. Yap, N. D. Reeves, S. Rajbhandari, and J. Spragg, "Fully convolutional networks for diabetic foot ulcer segmentation", *Comput. Biol. Med.*, vol. 89, Oct. 2017, pp. 53–65.
- [26] Q. Abbas, M. E. Celebi, I. F. García, and M. Rashid, "Unified approach for lesion border detection in skin images", *Comput. Methods Programs Biomed.*, vol. 109, no. 3, Mar. 2013, pp. 201–213.
- [27] J. I. Orlando, E. Prokofyeva, and M. B. Blaschko, "A discriminatively trained conditional random field model for blood vessel segmentation", *IEEE Trans. Med. Imaging*, vol. 36, no. 11, Nov. 2017, pp. 2418–2428.
- [28] T. Zhou, S. Ruan, and S. Canu, "A review of deep learning for medical image segmentation," *Pattern Recognit.*, vol. 93, Sept. 2019, pp. 123–134.
- [29] M. Goyal, M. H. Yap, N. D. Reeves, S. Rajbhandari, and J. Spragg, "Diabetic foot ulcer image analysis using deep learning", *IEEE Trans. Med. Imaging*, vol. 38, no. 3, Mar. 2019, pp. 652–661.
- [30] C. Wang, W. Yan, M. Smith, T. Filleter, and K. Najarian, "Deep learning-based diabetic foot ulcer detection", *Comput. Biol. Med.*, vol. 124, Art. no. 103939, Sept. 2020.
- [31] C. Shorten and T. M. Khoshgoftaar, "Image data augmentation for deep learning: A survey", *J. Big Data*, vol. 6, no. 1, Art. no. 60, July 2019.
- [32] L. Perez and J. Wang, "Effectiveness of data augmentation in CNN-based image classification", *Neural Comput. Appl.*, vol. 30, no. 7, Apr. 2018, pp. 1–13.
- [33] JayaLakshmi, G., Madhuri, A., Vasudevan, D., Thati, B., Sirisha, U., & Surapaneni, P. P. (2023). Effective disaster management through transformer-based multimodal tweet classification. *Revue d'Intelligence Artificielle*, 37(5), 1263.
- [34] S. Ruder, "Multi-task learning in deep neural networks: A survey", *Neural Netw.*, vol. 109, Jan. 2019, pp. 1–15.
- [35] Y. Zhang and Q. Yang, "Multi-task learning survey," *IEEE Trans. Knowl. Data Eng.*, vol. 34, no. 12, Dec. 2021, pp. 5586–5609.
- [36] S. Sun, Z. Cao, H. Zhu, and J. Zhao, "Optimization methods in machine learning", *IEEE Trans. Cybern.*, vol. 50, no. 8, Aug. 2020, pp. 3668–3681.
- [37] G. Litjens et al., "A survey on deep learning in medical image analysis", *Med. Image Anal.*, vol. 42, Dec. 2017, pp. 60–88.
- [38] M. T. Islam, M. Goyal, N. D. Reeves, and M. H. Yap, "Ensemble deep learning for diabetic foot ulcer diagnosis", *Expert Syst. Appl.*, vol. 232, Art. no. 120741, Jan. 2024.
- [39] C. Huang, Y. Zhang, and X. Li, "Multi-task diabetic foot ulcer segmentation and severity estimation", *IEEE J. Biomed. Health Inform.*, vol. 27, no. 6, June 2023, pp. 2791–2802.

- [40] S. Wang, M. Zhou, Z. Liu, Z. Liu, D. Gu, Y. Zang, and D. Dong, "Central focused convolutional neural networks for medical image segmentation", Med. Image Anal., vol. 59, Art. no. 101570, Jan. 2020.