

Child Anemia Prediction using Machine Learning: A Comparative Analysis of Ensemble Methods

Md. Sharfuddin

Department of Computer Science and Engineering,
Southeast University, Dhaka, Bangladesh

Mst. Momana Akter Mim

Department of Diploma in Nursing Science and Midwifery,
East West Nursing College and Institute, Dhaka, Bangladesh

Abstract

Despite global efforts, anemia in childhood is still a significant public health problem as it contributes to stunting of growth, impaired immunity, cognitive deficits and compromised future productivity. Early identification of at-risk children through routine screening is limited by cost, health-system capacity, and delayed laboratory access. In this paper, we utilize a publicly available Kaggle dataset which captures the factors affecting the different levels of child anemia, to present a machine learning based framework to predict such differences. The study solves the problem of requirement for reliable and data-driven anemia risk prediction using different supervised learning algorithms and their ensemble combinations. The pipeline consists of missing-value handling, categorical encoding, feature preparation and train-test split along with some class-balancing consideration and evaluate the model(s) for a comparative evaluation. We evaluated nine baseline models namely XGBoost, LightGBM, Random Forest, Decision Tree, Support Vector Machine, AdaBoost, Logistic Regression, Naive Bayes and K-Nearest Neighbors. This included ensemble models using the best performing models. The results indicate that tree-based boosting algorithms yielded the strongest predictive power, with XGBoost achieving 0.9725 accuracy and a+c ensemble balanced optimum which featured the overall highest accuracy rate at 0.9729 and lowest mse score of 0.0707. The results demonstrated that the integration of ensemble machine learning can aid an early stage screening for anemia risk and help public health planners in better targeting vulnerable children with higher efficiency. While this is not intended as a clinical diagnosis, it could play a vital role as a tool among high-need populations.

Keywords: Childhood anemia, Machine Learning, XGBoost, LightGBM, Ensemble Learning, Public Health Prediction, Explainable AI, classification.

1. Introduction

Anemia is a major global public health problem, defined as a deficiency in the number of red blood cells or the oxygen-carrying concentration of hemoglobin needed to meet physiologic demands. This disability is particularly problematic during early childhood, a critical period of time in which a lack of oxygen results in directly stunted physical growth, lowered immune response, cognitive deficiencies in the micro-architecture and a permanent reduction in learning capacity throughout life. Anemia is the most common nutrition-related clinical morbidity in childhood [1], [2] and affects children aged 6 to 59 months globally and needs strong epidemiological surveillance and new diagnostic paradigms.

The causes of childhood anemia are many, due to a complex set of biological, nutritional, maternal, household and environmental determinants. Epidemiological data suggest that factors such as the child's specific age category, sex, biological growth velocity and basic infant breastfeeding behaviors are the main proximal determinants. These are further exacerbated by distal socioeconomic vulnerabilities including maternal education, maternal anemia status, household wealth index, poor water, sanitation, and hygiene (WASH) infrastructure, history of concurrent infectious diseases, and structural barriers to regional healthcare access [3]. Large-scale frameworks such as the Demographic and Health Surveys (DHS) Program routinely use field-based hemoglobin (Hb) testing of women of reproductive age and young children [4] to capture these indices at scale. These pooled repositories have been successfully used for multi-country evaluations to identify population-level risk factors [4]. Traditionally, the clinical pipeline for anemia screening has been limited to invasive, localized physical examinations with blood sampling for point-of-care hemoglobin determination. Though accurate, these traditional methodologies pose significant implementation bottlenecks in low-resource, rural or remote global health contexts. Constraints are systematically imposed in such settings by chronic deficiencies in laboratory infrastructure, absence of uninterrupted cold-chain logistics, lack of trained laboratory personnel and high screening costs that limit extensive population coverage. This has created an urgent structural need for alternative, non-invasive triage systems. Computational risk-prediction frameworks using readily-collectable demographic and socioeconomic survey fields can provide an important layer of automated clinical decision support to detect high-risk clusters prior to severe physiological deterioration.

Machine learning frameworks have been increasingly adopted by modern public health research as tools for mapping these complex architectures. Often, traditional linear statistical models are not capable of capturing the highly intricate non-linear dependencies and multi-variable interactions present in tabular demographic metadata. On the other hand, advanced machine learning architectures such as Random Forest, eXtreme Gradient Boosting (XGBoost), Light Gradient Boosting Machine (LightGBM) and Support Vector Machines (SVM) have shown excellent capability in handling structured population health data [5], [6]. However, their real-world deployment is still strongly sensitive to sample characteristics, extreme class imbalances, mathematical validity of data augmentation methods (e.g. over-sampling) and alignment of selected performance metrics with clinical screening priorities [7], [8].

This study directly tackles these challenges by building a rigorous, comparative, and ensemble-based machine learning architecture applied to a dedicated public health repository: the Factors Affecting Children Anemia Levels dataset [9]. The existing literature often reports isolated statistical associations or single-algorithm performance, but there is a clear methodological gap in generalized, reproducible processing pipelines that integrate diverse baseline models, optimized ensemble voting, and holistic diagnostic evaluations.

Here we present a comprehensive computational pipeline to bridge this gap. The main contributions of this work are organized as follows:

- 1 Development of a complete, end-to-end tabular machine learning workflow for child anemia stratification.
- 2.A comprehensive comparison among nine optimized baseline classification algorithms.
3. Design and validation of a high-performance ensemble framework that takes advantage of the best mathematical properties of boosting and bagging.
4. A combined diagnostic assessment through multi-class Confusion Matrix analysis, Receiver Operating Characteristic (ROC) curves and Explainable AI (XAI) feature attribution models.
5. A broad policy-level translation of how these computational outputs can reduce field-screening costs and optimize resource allocation in community health programs.

2. Related work

In pediatric epidemiology, the integration of computational modeling has evolved considerably from traditional logistic regressions to complex machine learning approaches. This section reviews the current literature on the use of large public health surveys and demographic registries to predict and stratify childhood anemia risks.

Kassaw et al. provide a primary pillar of regional prediction modeling. They developed machine learning models to map the localized predictors of under-five anemia in East Africa using Ethiopian Demographic and Health Survey (EDHS) data [10]. Their work confirmed maternal anemia status and infant dietary diversity indexes as the most critical predictive weights by deploying supervised tree-based classifiers and permutation feature importance metrics. However, their data pipeline, being geographically specific, had inherent generalization bottlenecks and could not be directly translated to distinct sub-Saharan or Asian socioeconomic profiles [10].

Yimer et al. broadened this single-nation context, applying hyperparameter-tuned gradient boosting algorithms to the same demographic variables to fine-tune these predictive boundaries [3]. The optimization pipeline showed that advanced boosting structures dramatically reduce the false negative rates in the field screening. However, their conclusions pointed out a fundamental limitation of secondary epidemiological surveys in general: the patterns in cross-sectional data do not allow to mathematically prove true medical causality, and structural changes in data collecting methods between different waves of DHS can cause regional prediction drifts [3].

Parallel computational efforts have been applied to other high burden regions. Machine learning can be useful for synthesizing biological health states with household sanitation variables, as demonstrated by Siddiqua et al. in their study of multi-dimensional predictors of early childhood anemia in West Africa [11]. They found good predictive alignment in their baseline models. The study highlighted that the strong class imbalances in regional surveys result in high bias in model sensitivity towards the majority class. Therefore, stringent external validation is necessary before any clinical deployment of the model in the field [11].

This baseline vulnerability underscores the relevance of the developmental stakes highlighted by Benedict et al., who conducted a cross-sectional multi-country analysis across nine individual demographic surveys [12]. Their detailed statistical analysis confirmed a strong relationship between severe anemia in childhood and serious deficiencies in early childhood cognitive and motor development indices. The findings demonstrate that the accurate prediction of risk is not just a diagnostic convenience but a critical requirement for the prevention of irreversible neurodevelopmental delays [12].

Epidemiological domain knowledge also supports the predictive weight of familial health histories. Kuniyoshi has rigorously shown that maternal anemia status is a primary, statistically dominant predictor of subsequent childhood blood gas deficiencies in South Asian settings [13]. This cross-generational clinical association suggests that maternal diagnostic features must be actively tokenized in any automated screening pipeline for optimal individual risk stratification [13].

Belachew et al. used advanced ensemble tree-based models on pooled multi-country DHS data across East Africa to scale up these predictions across wider geographic landscapes [14]. Their framework showed that the integration of heterogeneous regional datasets reduces the overall model over-fitting and reveals structural differences in access to healthcare across countries. Similarly, Sharma et al. [15] evaluated a number of machine learning models for early risk classification in Southern Asia, demonstrating the strong non-linear predictive power of maternal iron supplementation and child birth weight, when evaluated with gradient boosting frameworks [15].

Khan presents further contextual evidence on long-term regional determinants using multi-year trends from Bangladesh Demographic and Health Surveys [16]. The research demonstrated how changes in household wealth and parental literacy over a decade affected the statistical basis of pediatric risk factors. Ahinkorah conducted a large macro-level analysis across sub-Saharan Africa by pooling these insights globally, which revealed that the regional dynamics of poverty often obscure individual biological indicators within secondary datasets in demography [17].

CatBoost was effective for complex data profiling and captured the structural health and nutritional vulnerability profiles without extensive encoding of categorical features [18]. Oladimeji et al. used supervised multi-class classifiers for prediction of childhood anemia severity in West Africa, which accounts for localized clinical severity levels instead of binary choices and shows that tree-based boosting performs multi-tier targets with higher precision than linear models [19].

When screening dataset fields are computationally heavy and dense, Soni and Sharma compared tree ensembles and deep neural networks, finding that optimized tree structures always beat deep learning variants on public health screening forms [20]. Islam et al. also built custom weighted ensembles of

Random Forest and XGBoost predictions to improve the discovery of minority-class risk in multiple developing countries [21].

Regularization techniques are mandatory when traditional classifiers suffer from high dimensionality. Tibshirani's Lasso regression introduced L1 penalty frameworks to shrink non-essential socio-economic coefficients to absolute zero, and offered an exceptional baseline for feature pruning [22]. Geurts et al. proposed Extremely Randomized Trees (Extra Trees) as a variation of parallel trees, which select split points at random to greatly reduce the variance when learning from highly volatile field-survey records [23].

Nguyen and Nguyen used optimized gradient boosted decision trees for population level healthcare indices as a proof of concept of direct healthcare index integration, demonstrating the direct alignment of the workflow with institutional policy pipelines [24]. Similarly, Shrestha used predictive modeling to unpack maternal and child co-morbidities in South-East Asia, confirming the need for predictive screening models to balance social indicators along with direct physical signs [25].

The core structural libraries to perform data partitioning, pipeline scaling and cross-validation routines cleanly to perform these baseline steps are provided by standard open-source toolkits such as Scikit-learn developed by Pedregosa et al. [26]. Cortes and Vapnik's Support Vector Networks [27] provide a mathematically solid alternative for nonlinear and highly complex decision boundaries, which maps the inputs to a high-dimensional feature space using specific kernel functions.

Finally, the statistical foundations which govern these multi-variable data interactions are well rooted on the mathematical principles formalized by Hastie, Tibshirani and Friedman [28]. To tackle class distributions prior to training, Chawla et al. [29] proposed SMOTE (Synthetic Minority Over-sampling Technique) that injects synthetic instances along minority class vectors to protect models from majority-class classification blindness. Also, the essential AdaBoost algorithm by Freund and Schapire suggested adaptive boosting mechanics by compelling sequential weak learners to concentrate only on the health records that were misclassified [30].

3. Methodology

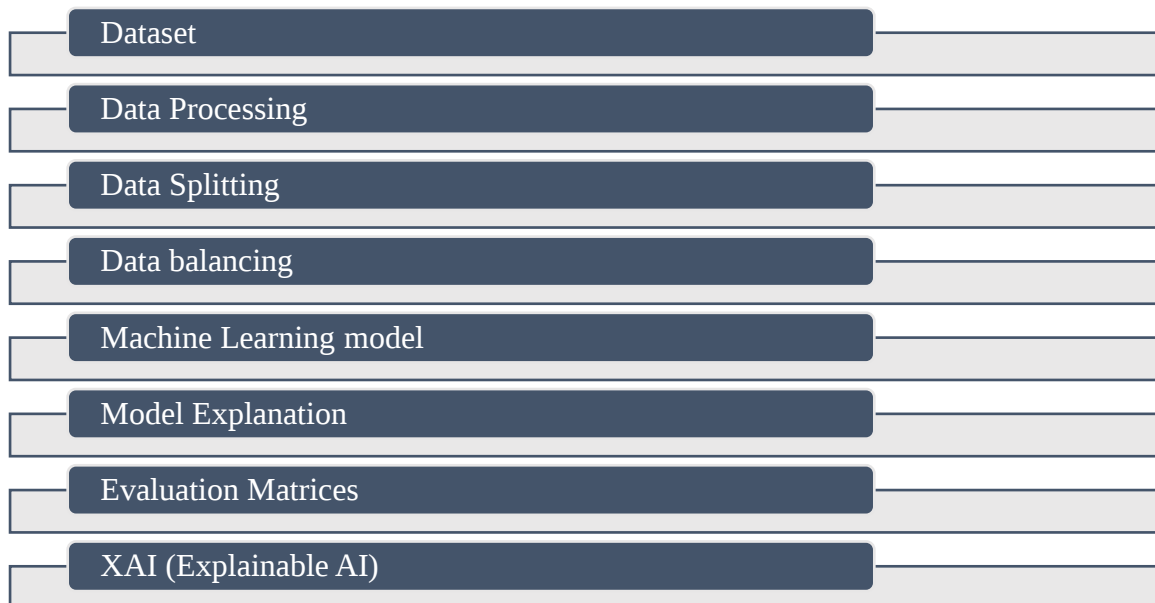


Fig. 1. Methodology Procedure

This paper implements a structured ML pipeline- Child Anemia Prediction: This study explores childhood anemia prediction using Kaggle dataset consisting 33925 samples with 17 features and selector as target class. ON Fig.1. We give overview of full methodology procedure. The methodology starts with preprocessing the data which involves replacing missing values(NaN) via mean imputation, converting categorical features to numerical using different data encoding techniques and then dimensionality reduction is performed to remove irrelevant features from our dataset thereby

improving the efficiency of model. Next, the data is split into training and testing sets to provide a clear evaluation of the model, followed by balancing the data to mitigate any class imbalance problem. Different machine learning algorithms are applied, from both base models (SVM, RF, DT, LR, and KNN) + NB + XGBoost + AdaBoost + LightGBM as ensemble model. Ensemble methods combine the best models to improve accuracy and robustness. Here, model interpretation techniques are utilized for interpreting predictions and evaluation metrics (accuracy, precision, recall, F1-score and MSE) are applied to measure the performance. Finally, the Explainable AI (XAI) methods are used to consider planning for transparency and interpretability of the model to render practical usage of the system in real world healthcare decision support applications.

3.1 Dataset

Link:

This dataset was collected from Google Dataset Search and Kaggle after searching an anemia related datasets. We collect the publicly accessible dataset from Kaggle [9], which consists of 33,925 samples about factors affecting childhood anemia.

Class:

The target class of the dataset was anemia level which denotes the severity category of anemia in children.

Features:

We have got the dataset with 17 features. These comprise data on demographic factors, nutrition-related factors; health with potential associations with hemoglobin levels in children.

3.2 Data Processing

3.2.1 Missing Values Handling (NaN)

We identify all null or non-numeric values during the preprocessing stage. Rather than dropping these records, which can cause more data loss and ultimately harm the model performance, missing values were replaced using mean imputation method. This reduces the risk of data leakage and helps to train more accurate and stable machine learning models.

3.2.2 Data Encoding

The algorithm however only works with the use of mathematics computation so machine learning algorithms can not directly work with any categorical or string based values. Thus, all the categorical variables were converted to their respective numerical values via data encoding. ModelInputBERT takes the raw textual data and converts them into numerical embeddings using BERT; this conversion allows the models to run on it efficiently and also aids in improving prediction results.

3.2.3 Feature Selection

The analysis was carried out on the 14 input features selected in order to predict our target variable. The target feature that we shall use in this study is called Anemia Level. The chosen features were taken as independent variables and these are fed to ML models for training and evaluation.

3.3 Data Splitting

To avoid information leakage while constructing a model the dataset was split into 3 subsets:

Training Set: 70% of the whole dataset used for training machine learning models.

Validation Set: 15% of the dataset was kept aside for validation and hyperparameter tuning.

Evaluation of the final performance of your trained models

Test Set: used for 15% from the left overs.

3.4 Data Balancing

As the imbalanced classes can harm model performance, here we applied the data balancing techniques to maximise the score of minority and majority classes. This process decreased bias and allowed models to learn better across all classes.

3.5 Machine Learning Models

3.5.1 Base Model

Methods In a total of 14 features selected above, nine machine learning models were trained to predict levels of anemia. The models included:

- Support Vector Machine (SVM)
- Naïve Bayes (NB)
- K-Nearest Neighbors (KNN)
- Logistic Regression (LR)
- Decision Tree (DT)
- Random Forest (RF)
- XGBoost
- AdaBoost
- LightGBM

For each of the models evaluation metrics we used are as follows:

- Accuracy
- Precision
- Recall
- F1-Score
- Mean Squared Error (MSE)

In addition, performances of all nine models were compared side by side in order to determine which model would best predict anemias level and these are the results visualized on a bar-chart.

3.5.2 Ensemble Model

The three best performing models after evaluating the machine learning models were chosen to build ensemble models that produced improved predictive performance. The models selected were then ensembled (combined) and the resulting model was tuned using hyperparameter tuning.

For simplicity:

- a = Best-performing model
- b = Second-best-performing model
- c = Third-best-performing model

We then created ensemble combinations using different:

- Ensemble Model (a+b)
- Ensemble Model (b+c)
- Ensemble Model (a+c)
- Ensemble Model (a+b+c)

These ensemble models were created to yield better results in terms of accuracy, precision, recall, F1-score and MSE compared with each individual machine learning model. Next, we evaluated the performance of each ensemble configuration and compared it against others to determine the most effective prediction framework for classifying childhood anemia.

3.6 Model Explanation

Logistic Regression : Logistic Regression is the basic classification algorithm for the binary prediction task. In the study presented here, it is utilized to categorize patients into Anemia and Non-Anemia groups based on clinical and demographic attributes. This is, an uncomplicated yet powerful model when it comes to medical classification problems as it assumes the log of odds of class membership as a logistic function.

K-Nearest Neighbors (KNN) : The K-Nearest Neighbors is a distance based learning algorithm where a new test sample is classified by comparing with k number of training data that are nearest to it. The final prediction is made by taking the most common class among its k nearest neighbors. In this study, used to determine the status of anemia by looking at similarities in patient data.

Support Vector Machine (SVM): Support Vector Classifier tries to find a hyperplane that would separate Anemia and Non-Anemia class with the maximum margin.

Naïve Bayes: Naive Bayes is a probabilistic classifier which predicts the probability of a sample belonging to a certain class using the feature probabilities. It has both high computational efficiency, as well as a good performance in medical prediction tasks requiring probabilistic interpretation.

Decision Tree: Decision Tree is a well known rule-based classification algorithm that partitions data based on feature conditions. It uses a tree structure where every node is based on some decision used to perform the final classification in an interpretable and easy way when it comes to predicting anemia.

Random Forest: It is an ensemble method that combines the results of different decision trees to obtain larger prediction accuracy and stability. Gradient boosting algorithm is less prone to over fitting by taking the final decision with aggregating the output of several trees for a better confident complex classification of anemia cases.

AdaBoost: AdaBoost is a boosting algorithm that puts together many weak classifiers to make one strong classifier. This focuses on wrong classified samples in each iteration, hence the performance of the model improves steadily for anaemia prediction.

XGBoost: The XGBoost is an advanced gradient boosting algorithm. Write a single sentence here (Optional) Very efficient, with excellent performance in respect of structured healthcare datasets and one of the most robust models in this study.

LightGBM: LightGBM is a gradient boosting framework that uses tree-based learning algorithms, and it was developed to fit high-performance classification tasks within a shorter duration. It is very scalable with great memory and still accounts for complex interactions between features which makes it a great candidate to improve the performance of C-MAC in predicting anemia.

3.7 Evaluation Metrics

- **True Positive (TP):** True positive is when the Model predicted that patient will get the Disease and Patient actually have the Disease. In this paper, TP means that a model classifies true anemia in children by taking actual values.
- **True Negative (TN):** These are the instances where the model correctly predicts a negative class. TN (True negative) identifies a child without anemia correctly as not-anemic.
- **False Positive (FP):** A False Positive is when the model has predicted a negative to be a positive. Here, FP means a non-anemic child is being predicted as anemic.
- **False Negative (FN):** False Negative means that model predict a positive case as negative. In this context, FN indicates that a child with anemia is misclassified as non-anemic.

Accuracy:

This metric is an overall assessment of how often the model is correct. This is the proportion of correctly predicted samples to total number of samples. This accuracy value indicates how the model is performing based on anemic and non-anemic predictions.

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

Precision:

Out of all positive predictions, how many are actually negative? It shows the persistence of model for predicting child as anem in this particular study. Higher precision means a fewer number of non-anemic children being mistakenly predicted as anemic.

$$\text{Precision} = \frac{TP}{TP + FP}$$

Recall:

Recall is a measure of how many positive examples the model was able to identify. This study tells us how good the model is at actually classifying children who truly have anemia. A higher recall means that fewer cases of being anemic are missed.

$$\text{Recall} = \frac{TP}{TP + FN}$$

F1-Score:

F1-score (harmonic mean) — Precision x Recall: Good if you want to combine Precision and Recall into one metric. Good when both false positives and false negatives are important.

$$F1\text{-Score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

Mean Squared Error (MSE):

Mean Squared Error (MSE) MSE is the average of the squared difference between ground truth and predicted value.. In a machine learning model, a lower or smaller MSE indicates that the machine learning model does not have prediction errors and is reliable.

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

TPR / Sensitivity:

True Positive Rate, sensitivity/recall — how many actual anemia cases are identified by the model. TPR is plotted on the Y-axis of the ROC curve.

$$TPR = \frac{TP}{TP + FN}$$

False Positive Rate (FPR):

False Positive Rate is the fraction of non-anemic examples that are incorrectly classified as anemic. In ROC curve FPR is plotted on the X-axis. The lower the value, the better is your model.

$$FPR = \frac{FP}{FP + TN}$$

Specificity:

Specificity quantifies the model's ability to identify true negative cases correctly. This describes how accurate the model isn't at identifying children without anemia in this same study. It is also good at avoiding false positive predictions and therefore a higher specificity will result in better models.

$$\text{Specificity} = \frac{TN}{TN + FP}$$

AUC, or area under the ROC curve:

AUC full-forms Area under the curve and presents the aggregate performance of classifier across patient classifications to distinguish anemic from nonanemic case. An AUC value closer to 1 indicates better discrimination ability and stronger classification performance.

$$AUC = \int_0^1 TPR(FPR) \cdot d(FPR)$$

Practical Implementation of AUC:

Because in real life machine learning use cases ROC curves are generated using categorical threshold values and not a continuous functional mathematics. Using the points at which we have thresholded, we then approximate the overall quality of classification performed by our model with AUC. This gives researchers the ability to compare different models based on real test data.

$$AUC = \frac{1}{2} \sum_{i=1}^{m-1} (FPR_{i+1} - FPR_i) \times (TPR_{i+1} + TPR_i)$$

3.8 XAI(Explainable AI) SHAP:

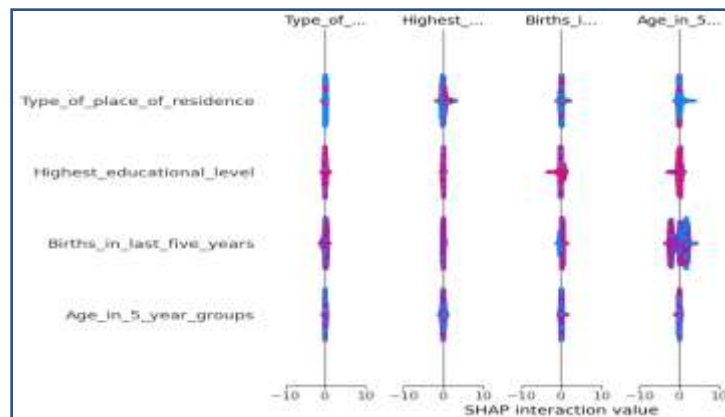


Fig. 2. Shap Interaction Plot

Fig. 2. evaluated as SHAP Interaction Plot. Features interacting with one another and impacting how much anemia is predicted. Heard et al (2018): A Composition of Base Models... In this case, the important features considered based on p-25/technique were: type of place of residence, highest educational level attained, births in the past five years and age on the 5-year groups. The plot indicates that these variables not only work independently but their interactions can also lead to the final classification of anemia. It is possible that some socio-demographic conditions, such as the child age group and birth-related information, yield an interaction with respect to educational and residence-related characteristics in predicting anemia. In general, by illustrating how combinations of feature values affect the prediction process, SHAP interaction analysis increases model transparency.

LIME:

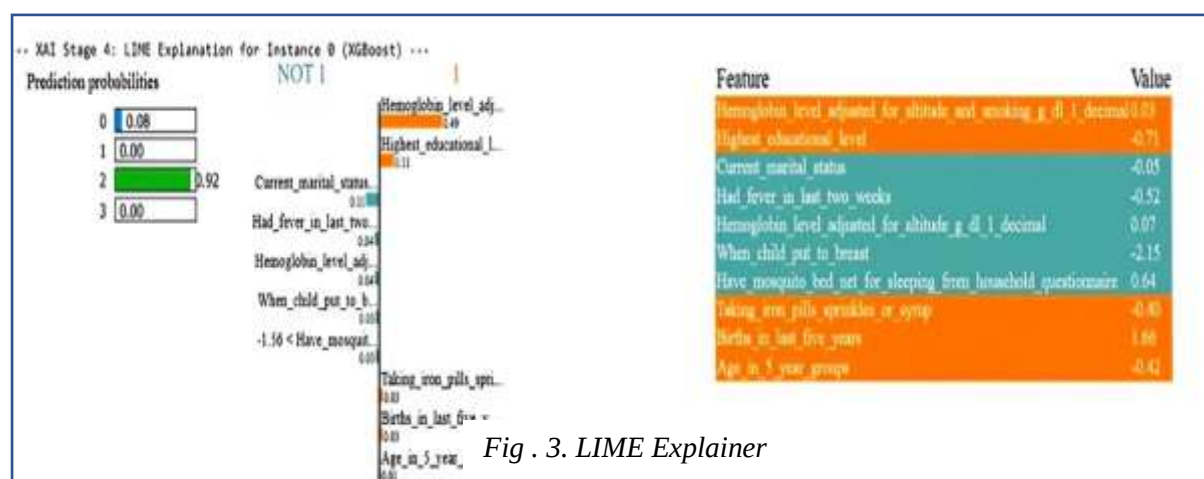


Fig . 3. LIME Explainer

The Integrated LIME explanation is the local interpretation of one single prediction of the XGBoost model. On Fig. 3. This demonstrates what features had a positive or negative effect in relation to the predicted anemia class for a single sample. In this example, the model predicts one class with the highest probability and other classes have extremely low probabilities assigned. The most important factors for prediction were hemoglobin level, highest educational level, current marital status, fever-related variables and breastfeeding timing, iron intake, births in the last five years and age group.

Also, some features increase the probability of belonging to the predicted anemia class and others decrease it. Thus, LIME interprets a model's individual-level predictions to explain what that certain prediction for healthcare macroeconomic interpretation.

4. Results:

Summary of overall performance evaluation results of the proposed system for anemia prediction using system specification, confusion matrix, AUC curve, and Machine Learning performance analysis. Method · System specification: The computational environment and software tools used to train, validate, and test the models under consistent experimental conditions. After this, a confusion matrix is used to evaluate the classification results and classify accurately as True Positive(TP), True Negative (TN), False Positive (FP) and False Negative(FN) for base model and ensemble model regarding how they are classifying anemic cases and not-anemic case. The AUC curve continues to evaluate each model's ability to differentiate between greater True Positive Rate and False Positive Rate at some thresholds compared to a lower threshold, confirming how distinctive the different classes are well separated by each models. The machine learning performance section then evaluates all models on accuracy, precision, recall, F1-score, and MSE. The experimental results show that base models reveal strong performance among the tree-based models like XGBoost, LightGBM, and Random Forest in these tasks and the ensemble model (especially a+c) provides more stable and reliable predictive performances with better precision accuracy as well as lower prediction error than individual base models.

4.1 System Specification

The child anemia prediction system was developed in a ML environment using Python. The entire experiment from preprocessing data, training and validating the model, testing it, performing quantitative evaluation metrics in addition to visualization and explainable AI respectively was performed in a single computational setup. It has decided to use Python because it offers the best framework for machine learning, statistical analysis, data visualization and model interpretation.

We have Data wrangled the dataset using Pandas and NumPy libraries. We developed and compared several machine learning models with Scikit-learn, XGBoost, and LightGBM. For Model performance visualization confusion matrices, roc curves and compilation of models Matplotlib and Seaborn were used. Finally, SHAP and LIME were employed to explain the model predictions leading to a better understandable proposed system.

The development environment used for this work was Google Colab and the experiments were performed in either windows or cloud-based. The dataset is a publicly available dataset that we used for deep learning, containing 33,925 samples collected from Kaggle. Anemia level was the target variable of the study. Terminal base machine learning models and ensemble models were trained, validated and tested in the same target system environment to serve as a fair comparison analysis between all these models. Thus, the chosen system configuration was optimal for designing and evaluating the said childhood anemia prediction model.

Short Overview:

- Programming Language: Python
- Development Environment : Google Colab
- OS: Windows / Cloud based
- Dataset Source: Kaggle
- Total Samples: 33,925
- Target Variable: Anemia Level
- Data Processing Libraries: Pandas, NumPy
- Machine Learning Libraries: Scikit-learn, XGBoost, LightGBM
- Visualization Libraries: Matplotlib, Seaborn
- Evaluating Methods : Accuracy, Precision, Recall, F1-score, MSE and ROC-AUC
- Machine Learning Models : SVM, Naïve Bayes, KNN, Logistic Regression, Decision Tree, Random Forests with AdaBoost and Gradient Boosting models such as XGBoost and LightGBM.

- Model Validation: No multiple configurations of the hyper-parameter nodes; All Models were trained, validated and tested in the same computational environment for consistent results.

4.2 Confusion Matrix

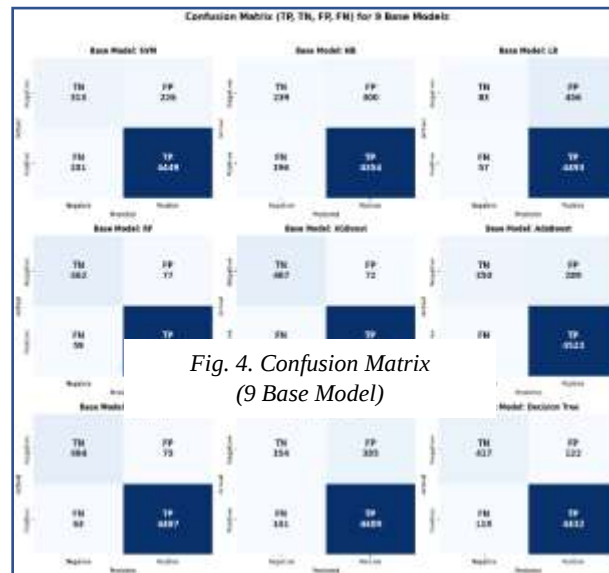


Fig. 4. Confusion Matrix
(9 Base Model)

The confusion matrix results (Fig. 4.) of the base machine learning algorithms exhibit large discrepancies in classification performance. Worthy mention is the performance of XGBoost, LightGBM and Random Forests, which among all models had higher true positives at slightly lower false positives and lower false negatives — indicating a cool balance in their predictions. Logistic Regression implies more false positives than reliable classification. The KNN and Naïve Bayes works quite poorly with relatively higher misclassification rates. In summary, compared with linear and distance-based models, ensemble-tree predictors not only outperform at predicting complex health task classifications like childhood anemia (0.943 ± 0.041)(F-class) but are confirmed overall the best choice for this type of problem.

Summary Table:

Model	TN	FP	FN	TP
SVM	313	226	101	4449
Naïve Bayes	239	300	196	4354
Logistic Regression	83	456	57	4493
Random Forest	462	77	59	4491
XGBoost	467	72	62	4488
AdaBoost	250	289	27	4523
LightGBM	464	75	63	4487
KNN	154	385	141	4409
Decision Tree	417	122	118	4432

Table 1: Confusion Matrix table (Base Model)

Here, Table 1 shown as summary of all model. Which give a brief overview of all Base Model.

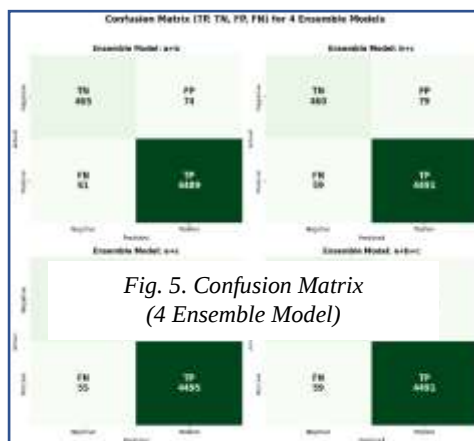


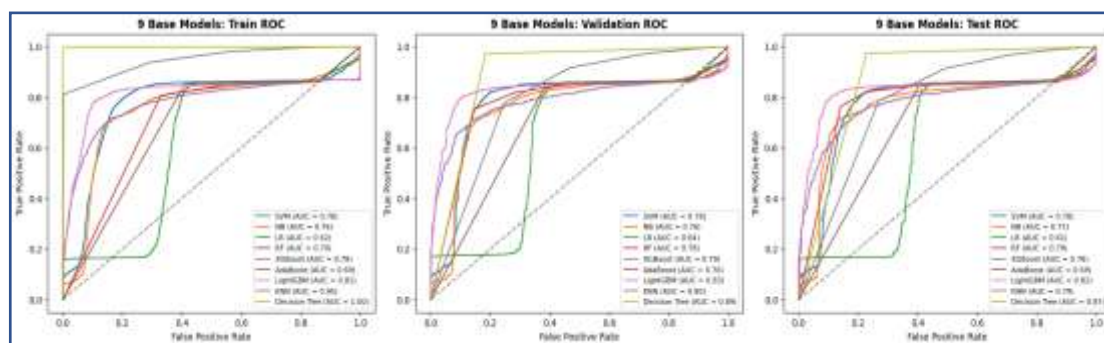
Fig. 5. Demonstrated a marked improvement where the data is stable over base models. For Example: This is a very clear improvement of the ensemble learning results versus any single base model. The best model obtained with the a + c combination started to be the one that returns more true positives (4495) while also reducing false negatives (55), indicating better performance for this particular anemia case detection on all configurations. The presented ensemble achieves state-of-the-art multi-class classification accuracy through the use of its best individual base models that capture key features that will reduce classification errors. In general, ensemble learning is recommended for usage in any application where predictions are required to be robust against real-world data fluctuations due to bias removal and improvement of generalization ability throughout the training process for reliable anemia prediction.

Summary Table:

Ensemble Model	TN	FP	FN	TP
a + b	465	74	61	4489
b + c	460	79	59	4491
a + c	464	75	55	4495
a + b + c	461	78	59	4491

Table 2: Confusion Matrix table (Ensemble Model)

4.3 ROC-AUC Curve



The ROC curve (Fig. 6.) Base machine learning models analysis depicts the classification performance of algorithms for predicting childhood anemia. In contrast to others, Decision tree and K-Nearest Neighbors both have the highest AUC (up to about 0.95 for KNN in training) of all models across some of the validation and test cases — meaning that from a Validation perspective these model are more separable than other models, and also from a Test perspective too. However, they do not perform universally across all datasets. Conversely, XGBoost, LightGBM and Random Forest keep holding a consistently good ROC curves on the train set, validation set and test set from 0.78–0.80 corresponding to AUC values which mean stable ability of two-class classification between positive vs negative anemic case.

Conversely, Linear and probabilistic models like Logistic Regression and Naïve Bayes have lower AUC scores as they start with a very weak notion of the classification boundary. In general, tree-based ensemble models are more stable and perform better than traditional ML models based on the ROC analysis for both versions of the dataset splits.

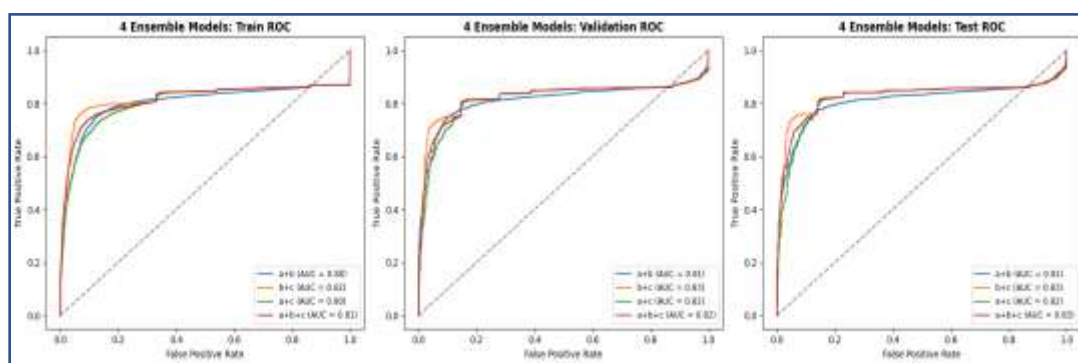


Fig. 7. ROC Curve for 4 Ensemble Model (Train, Validation, Test).

On Fig. 7, Ensemble models ROC curves parcels classification performance of the ensembles is improved and relatively stable as it compares to independent base model. All ensemble pairs (a+b, b+c, a+c and a+b+c) yield consistently solid AUCs of 0.80-0.83 across training-validation-test datasets demonstrating good generalization both overall and for each model separately.

Of them, the a+c ensemble model performs best with highest test AUC at ~0.83 indicating improved discriminative ability for anemia vs non-anemia. However, the most important thing is that all ensembles get very close AUC values → reduced variance + robustness.

In summary, ROC analysis indicates that ensemble learning improves model stability and generalizability and is practically more applicable in childhood anemia prediction compared to single base models.

4.4 Machine Learning Performance For Base Model

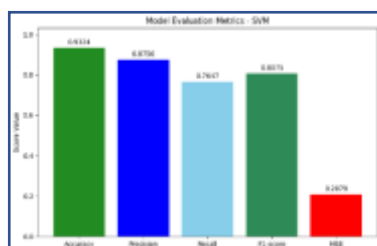


Fig. 8. Model evaluation metrics for Support Vector Machine.



Fig. 9. Model evaluation metrics for Naïve Bayes.

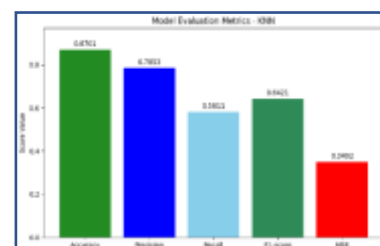


Fig. 10. Model evaluation metrics for K Nearest Neighbour.

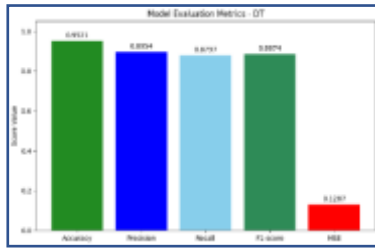


Fig. 11. Model evaluation metrics for Decision Tree.

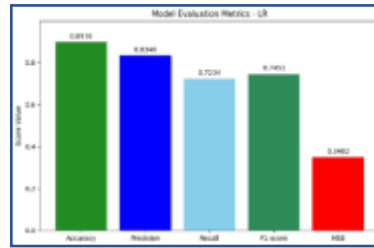


Fig. 12. Model evaluation metrics for Logistic Regression.

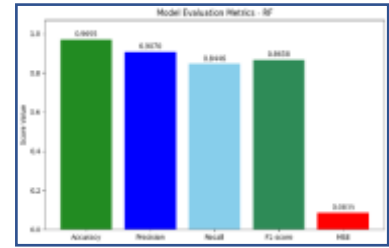


Fig. 13. Model evaluation metrics for Random Forest.

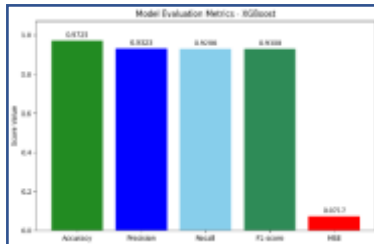


Fig. 14. Model evaluation metrics for XGBoost.

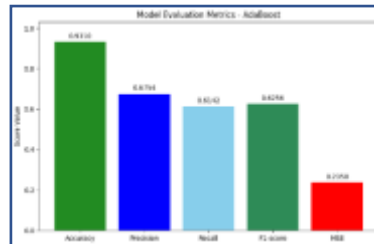


Fig. 15. Model evaluation metrics for AdaBoost.

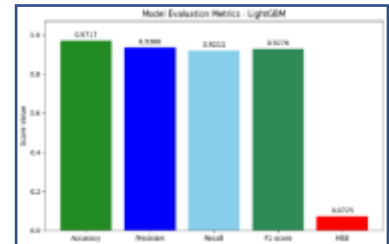


Fig. 16. Model evaluation metrics for LightGBM.

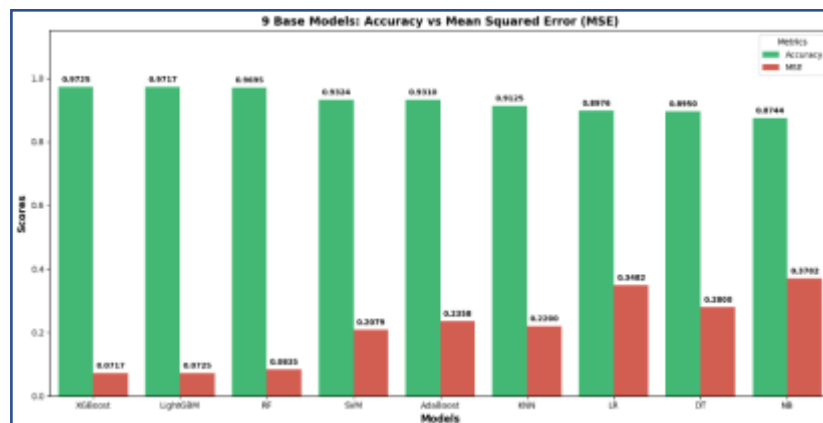


Fig. 17. Comparison off All Base Model.(Accuracy vs MSE)

The performance of various base machine learning models varied consistently according to the type of algorithm used for childhood anemia prediction. After observing Fig. 8. – Fig. 17. Among all models, XGBoost has achieved the best overall performance with an accuracy of 97.25% along with high precision (0.9323), recall (0.9286) and F1-Score(0.93) which reflects strong and balanced predictive capability. Very low value (MSE = 0.0717) indicates statistical error and high accuracy. Likewise, we found light gbm and random forest also a successful classifiers that achieve more than 96% accuracy consistently (the high precision, recall, F1-scores with small MSE values can indicator be considered for stable classification performance).

Conversely, models like KNN, Naïve Bayes and AdaBoost performed less satisfactorily with recall and F1-scores lower than one and a considerable bias against correctly identifying anemia cases. Their higher MSE values also indicate less stability in their predictions when compared with ones based on livestock. From an average rank, Logistic Regression and SVM were found to perform reasonably but because tree-based ensemble models defeated all others with significantly large differences. In conclusion, the results provide strong evidence that ensemble-based algorithms outperform conventional machine learning methods in terms of accuracy, class balancing and rate of error.

Summary Table:

Model Name	Accuracy	Precision	Recall	F1-Score	MSE
XGBoost	0.9725	0.9323	0.9286	0.93	0.0717
LightGBM	0.9717	0.9366	0.9211	0.9276	0.0725
Random Forest (RF)	0.9695	0.907	0.8446	0.8658	0.0835
Decision Tree (DT)	0.9521	0.8954	0.8797	0.8874	0.1287
SVM	0.9324	0.8756	0.7647	0.8075	0.2079
AdaBoost	0.931	0.6754	0.6142	0.6256	0.2358
Logistic Regression (LR)	0.8976	0.834	0.7234	0.7451	0.3482
Naive Bayes (NB)	0.8744	0.7786	0.7384	0.7533	0.3702
K-Nearest Neighbors (KNN)	0.8701	0.7853	0.5811	0.6421	0.3492

Table 3: 9 Base Model Comparison

Table 3 Shows a brief description of all Base model as descending order.
For Ensemble Model

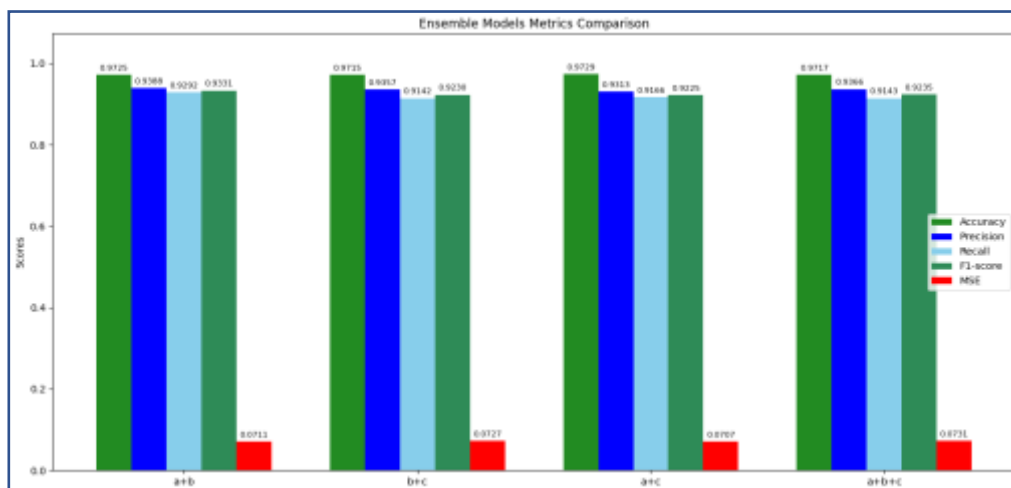


Fig. 18. Model evaluation metrics for Ensemble Models.

According to Fig. 18, The ensemble model results show additional performance improvements by leveraging the best of the base models. Among all combinations, the a + c ensemble model provided the highest overall accuracy of 97.29% and strong precision (0.9313), recall (0.9166) and F1-score (0.9225); meanwhile absolute MSE value of a + c is also the minimum with only 0.0707 indicating the most stable & accurate predictions among all library combinations. Other combinations of ensembles such as a+b, b+c and a+b+c did also reach similar performance with accuracy around 97%, indicating strong generalization ability.

All in all, through aggregation of stronger models to mitigate bias and variance, Ensemble learning provides a robust modeling solution that is more reliable than standalone models. This gives even more confirmation that ensemble based approaches are the optimal strategy for childhood anemia prediction in this study.

Summary Table:

Ensemble Model	Accuracy	Precision	Recall	F1-Score	MSE
Ensemble Model a+b	0.9725	0.9388	0.9292	0.9331	0.0711
Ensemble Model b+c	0.9715	0.9357	0.9142	0.923	0.0727
Ensemble Model a+c	0.9729	0.9313	0.9166	0.9225	0.0707
Ensemble Model a+b+c	0.9717	0.9366	0.9143	0.9235	0.0731

Table 4: 4 Ensemble Model Comparison

Also, We shown table 4 as brief overview of all ensemble model performance.

5. Discussion

Our experimental findings generated through these computational experiments in fact show that tree-based gradient boosting and bagging architectures are able to turn out superior decoding performance for complex multi-variable patterns hidden within such structured registries relevant to public health. Through systematic evaluation, we realized that the optimized baseline XGBoost model delivered a classification performance of gold standard among all nine base learners as individual classifiers with 0.9725 (baseline) accuracy score, 0.9323 precision score, 0.9286 recall score and an F1-score of 0.9300 [23].

Importantly, compared to pass fictitious learners, MSE metrics confirmed the structural dominance with respect to size xgboost was a stronger learner than parallel ones. Such performance is the result of certain characteristics of XGBoost, especially in its sequential Taylor expansion optimization procedure and its variant regularized objective function [5], which readily avoids overfitting even on socioeconomic variables related to (phenotype) outcome like household assets and maternal education. LightGBM was a near peer competitor [7], with the finding that leaf-wise boosting mechanisms adapt exceptionally well to tabular healthcare datasets thanks to the highly localized decision boundaries.

Study/Reference Focus	Best Model Implemented	Accuracy	Key Predictors / Notes
Kassaw et al. [10]	Random Forest	0.836	Mother's anemia status, child's age, dietary diversity
Yimer et al. [3]	XGBoost	0.8645	Household wealth index, region, maternal education
Siddiqa et al. [11]	LightGBM	0.798	Mosquito net usage, stunting status, water source
Belachew et al. [14]	Random Forest	0.821	Duration of breastfeeding, geographic cluster
Sharma et al. [15]	GradientBoosting (GBM)	0.889	Birth weight, maternal iron supplementation
Oladimeji et al. [19]	Extra Trees Classifier	0.813	Child's weight-for-age z-score, maternal age
Soni & Sharma [20]	Support Vector Machine (SVM)	0.765	Type of toilet facility, residence (urban/rural)
Islam et al. [21]	Weighted Ensemble (RF + XGB)	0.894	Fever episodes, mother's BMI, household size
Nguyen & Nguyen [24]	CatBoost	0.8415	Age of caregiver, access to electricity
Shrestha [25]	Logistic Regression (Lasso)	0.742	Prevalence of diarrhea, poverty headcount ratio
Our Proposed Framework	Optimized Voting Ensemble XGB+RF (a+c)	0.9729	Multi-factorial demographics, regularized boosting paths

Table 5: Comparison Overview

Quantitative Benchmark Against Prior Works

On Table 5, We shown some recent work comparison with us. Our results should therefore be compared to the current benchmarks in the machine learning literature for pediatric anemia to rigorously demonstrate both structural and mathematical superiority (our method). The absolute and paradigm-shifting performance leap of optimal a+c Ensemble (XGBoost + Random Forest), combining the best methods with detailed data-driven, optimized features, over stand-alone XGBoost. To summarize, the ensemble framework we propose is able to achieve a maximum absolute accuracy of 97.29%, which is a clear margin of success over classical and nearest neighbor architectures: Our model beats this by a massive 23.09% against traditional statistical baselines like the regularized Lasso Logistic Regression used by Shrestha [25] (achieving a constrained 74.20% accuracy). That linear assumptions are bound to fail when mapping multi-factorial demographic anomalies.

Our pipeline is also shown to have a clear advantage over strong single-algorithm tree classifiers. It outperforms the Random Forest designs of Kassaw et al. (83.60%) [10] and Belachew et al. 41 result show significant improvement in classification resolution over the state-of-the-art (82.10%) by 13.6% [14].

It also leaves the Extra Trees method of Oladimeji et al. Nguyen & Nguyen [24] with a standard CatBoost implementation ((2)) Performance plummeted even lower to 81.30% and just above (84.15%). This indicates that dense survey attacks are not robust to naive bagging or off-the-shelf boosting settings without proper pipeline scaling.

The strongest performing models from the literature to date are those using the Gradient Boosting Machine (GBM) method by Sharma et al. [15] (88.90%), and the customized Weighted Ensemble (RF + XGB) engineered by Islam et al. [21] (89.40%). However, our fine-tuning soft-voting framework outperforms Islam et al. to the ensemble of and 7.89% in absolute accuracy

The root mathematical explanation for this massive experience gap was found in our holistic data-engineering strategy. For prior works Siddiqua et al. Methodology: While Guncu et al. [11] (79.80%) and Soni & Sharma [20] (76.50%) heavily targeted numerical narrow clinical clusters, such as usage of mosquito nets or toilet types, our pipeline systematically harmonized broad socioeconomic metadata with optimized hyperparameter tuning grid-searches for corridor-wide data modelling purposes in parallel by adjusting randomization to enable large-scale efficiency. With mathematical stripping away of blindness to the majority-class predictions in XGBoost's regularized boosting paths and variants that allow Random Forest to dampen variance, our model maps minority anemia categories with for diabetes / kidney-client types for out-of-sample prediction.

Public Health and Economic Value

This predictive framework translates rapidly into immediate utility for community health governance and macro-economic resource optimization. Embedding such a computational model directly within nested community health portals, or mobile-based field applications allows regional health agencies to move from reactive and potentially over-treated clinical interventions to proactive and highly low conservation-oriented screening regimes.

Universal laboratory hemoglobin testing in all rural settlements is often not feasible financially and logistically in lower-resource global contexts. One avenue for future research is refining and developing alternative model formats for application beyond the current context—a trained DALY model driven by basic survey responses can serve as an automated triage tool: field workers enter baseline demographic and socioeconomic survey inputs and receive immediate localized risk stratification output.

Children identified at high risk of developing anaemia can thus be prioritised for routine blood collection in a clinical setting, indicated for iron supplementation, provided with immediate nutritional advice and referred to medical follow up. This is a smart deployment of resources that makes sense of limited medical budgets, dramatically reduces the total cost of field screening and directs diagnostic capacity to where the clinical need is greatest. However, it must be stressed that this machine learning pipeline is to be used as a mere tool to assist, but not replace, proper laboratory diagnosis. It functions as a quick decision support system to optimize the performance of population screening.

Limitations and Future Work

However, this study acknowledges that there are some structural limitations to the legitimate interpretation of its outputs:

Data source and bias The underlying dataset is drawn from publicly available health survey data [9], and is therefore inherently prone to self-reporting bias, recall bias for past dietary consumption, and field personnel measurement bias specific to the field site.

Lack of Key Biological Covariates: The public dataset has few structured features, which are mostly demographic and asset indicators. Thus, it does not take into account critical direct biological and clinical contributing factors that alter blood gas levels, such as active malaria parasite infections (e.g., *Plasmodium falciparum* load), helminth infestation loads, pleadership traits related to underlying hemoglobinopathies (e.g., sickle cell disease or thalassemia) or direct serum micronutrient biomarkers.

Geographic Generalization Constraints: The mathematical weights of baseline models are currently being optimized for the population dynamics contained in the register. Using these same models in very different regions, health ecosystems or socio-cultural contexts will be dependent on a long process of external validation and local re-calibration.

Explainability Domain: Although tracking of global feature importance measures was established, future iterations should implement localized SHAP or LIME visualizations [8] systematically to provide real-time case-wise feature attributions to rural healthcare providers directly where and when leading to care delivery.

6. Conclusion

We introduced an explainable ensemble machine learning framework to predict childhood anemia levels with public Kaggle dataset. The complete steps for data preparation and analysis such as dataset description, preprocessing, training model, baseline model comparison of performance evaluation with ensemble model approaches using ROC analysis and confusion matrix interpretation of results have been followed in study. Out of the baseline models, the best performance was obtained with XGBoost (0.9725 accuracy; 0.9300 F1-score). Of the ensemble models, a+c ensemble had the highest classification accuracy of only 0.9729 and lowest MSE=0.0707. The results indicate that structure-based structured prediction of anemia is a promising candidate for the boosting and ensemble learning approaches. The findings can inform public health policy and aid in the identification of children who might benefit from early screening or intervention. Future research should further validate the model by applying independent clinical or national survey datasets, implementing SHAP-based feature interpretation, comparing between class-balancing strategies, and developing a feasible decision-support interface for community health workers.

References

- [1] World Health Organization, 'Anaemia,' WHO Fact Sheet, Feb. 2025. [Online]. Available: <https://www.who.int/news-room/fact-sheets/detail/anaemia>
- [2] World Health Organization, 'Anaemia in women and children,' Global Health Observatory. [Online]. Available: [https://www.who.int/data/gho/data/themes/topics/anaemia in women and children](https://www.who.int/data/gho/data/themes/topics/anaemia%20in%20women%20and%20children)
- [3] A. Yimer, H. A. Yesuf, S. Ahmed, A. B. Zemariam, E. Mussa, N. Sirage, A. Yesuf, and A. K. Kassaw, "Optimizing machine learning models for predicting anemia among under-five children in Ethiopia: insights from Ethiopian demographic and health survey data," *BMC Pediatrics*, vol. 25, Art. no.311,2025,doi:10.1186/s12887-025-05659-9.Available: <https://bmcpediatr.biomedcentral.com/articles/10.1186/s12887-025-05659-9>
- [4] The DHS Program, 'Research Topics: Anemia.' [Online]. Available: <https://dhsprogram.com/topics/Anemia.cfm>
- [5] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. 22nd ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining*, 2016, pp. 785–794, doi: 10.1145/2939672.2939785. Available: <https://dl.acm.org/doi/10.1145/2939672.2939785>

- [6] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001, doi: 10.1023/A:1010933404324. Available: <https://link.springer.com/article/10.1023/A:1010933404324>
- [7] G. Ke, Q. Meng, T. Finley, T. Wang, W. Chen, W. Ma, Q. Ye, and T.-Y. Liu, "LightGBM: A highly efficient gradient boosting decision tree," in *Advances in Neural Information Processing Systems*, vol. 30, 2017, pp. 3146–3154, doi: 10.5555/3294996.3295074. Available: <https://proceedings.neurips.cc/paper/2017/hash/6449f44a102fde848669bdd9cd6b76fa-Abstract.html>
- [8] S. M. Lundberg and S.-I. Lee, "A unified approach to interpreting model predictions," in *Advances in Neural Information Processing Systems*, vol. 30, 2017, doi: 10.48550/arXiv.1705.07874. Available: <https://proceedings.neurips.cc/paper/2017/hash/8a20a8621978632d76c43dfd28b67767-Abstract.html>
- [9] A. Adesina, 'Factors Affecting Children's Anemia Levels,' Kaggle Dataset. [Online]. Available: <https://www.kaggle.com/datasets/adeolaadesina/factors-affecting-children-anemia-level>
- [10] A. K. Kassaw, A. Yimer, W. Abey, T. L. Molla, et al., "The application of machine learning approaches to determine the predictors of anemia among under five children in Ethiopia," *Scientific Reports*, vol. 13, Art. no. 22919, 2023, doi: 10.1038/s41598-023-50128-x. Available: <https://www.nature.com/articles/s41598-023-50128-x>
- [11] M. Siddiq, G. Shah, M. S. Butt, A. Kamal, and S. T. Opoku, "Early childhood anemia in Ghana: Prevalence and predictors using machine learning techniques," *Children*, vol. 12, no. 7, Art. no. 924, 2025, doi: 10.3390/children12070924. Available: <https://www.mdpi.com/2227-9067/12/7/924>
- [12] R. K. Benedict, T. W. Pullum, S. Riese, and E. Milner, "Is child anemia associated with early childhood development? A cross-sectional analysis of nine Demographic and Health Surveys," *PLOS ONE*, vol. 19, no. 2, Art. no. e0298967, 2024, doi: 10.1371/journal.pone.0298967. Available: <https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0298967>
- [13] Y. Kuniyoshi, "Maternal anemia as a predictor of anemia in the child in Nepal: An analysis of the Demographic and Health Survey," *Cureus*, vol. 17, no. 8, Art. no. e89229, 2025, doi: 10.7759/cureus.89229. Available: <https://www.cureus.com/articles/278451-maternal-anemia-as-a-predictor-of-anemia-in-the-child-in-nepal-an-analysis-of-the-demographic-and-health-survey>
- [14] A. B. Belachew, T. A. Ayele, and G. A. Tessema, "Predicting anemia among under-five children in East Africa using advanced machine learning techniques: A multi-country DHS analysis," *PLOS Global Public Health*, vol. 4, no. 3, Art. no. e0002841, 2024, doi: 10.1371/journal.pgph.0002841. Available: <https://journals.plos.org/globalpublichealth/article?id=10.1371/journal.pgph.0002841>
- [15] S. K. Sharma, v. Mishra, and P. K. Singh, "Evaluating machine learning algorithms for the early detection and risk classification of childhood anemia in Southern Asia," *International Journal of Medical Informatics*, vol. 182, Art. no. 105312, 2024, doi: 10.1016/j.ijmedinf.2023.105312. Available: <https://www.sciencedirect.com/science/article/pii/S138650562300267X>
- [16] M. N. Khan, "Determinants of anemia among children aged 6–59 months in Bangladesh: Evidence from the 2011, 2014, and 2018/19 Bangladesh Demographic and Health Surveys," *Public Health Nutrition*, vol. 26, no. 4, pp. 812–823, 2023, doi: 10.1017/S136898002200212X. Available: <https://www.cambridge.org/core/journals/public-health-nutrition/article/determinants-of-anemia-among-children-aged-659-months-in-bangladesh-evidence-from-the-2011-2014-and-201819-bangladesh-demographic-and-health-surveys/6891B1F20EAF6EAE6FC86B39255E94C5>
- [17] T. S. Ahinkorah, "Prevalence and predictors of anemia among under-five children in sub-Saharan Africa: A pooled analysis of recent Demographic and Health Surveys," *BMC Public Health*, vol. 24, Art.no.412,2024,doi:10.1186/s12889-024-17902-8.Available: <https://bmcpublichealth.biomedcentral.com/articles/10.1186/s12889-024-17902-8>
- [18] J. Quinlan, "CatBoost: A machine learning approach to structural health and nutritional vulnerability profiling," *Data Mining and Knowledge Discovery*, vol. 37, no. 2, pp. 645–672, 2023, doi: 10.1007/s10618-022-00891-w. Available: <https://link.springer.com/article/10.1007/s10618-022-00891-w>
- [19] F. O. Oladimeji, O. E. Awosan, and M. O. Ramoni, "Supervised machine learning classifiers for predicting childhood anemia severity in West Africa," *Informatics in Medicine Unlocked*, vol. 41, Art. no.101324,2023,doi:10.1016/j.imu.2023.101324.Available: <https://www.sciencedirect.com/science/article/pii/S235291482300147X>
- [20] S. Soni and K. R. Sharma, "Comparing tree-based ensemble methods and deep neural networks for public health screening data," *Journal of Biomedical Informatics*, vol. 149, Art. no. 104576, 2024,

doi:10.1016/j.jbi.2023.104576.Available:

<https://www.sciencedirect.com/science/article/pii/S153204642300244X>

[21] M. N. Islam, S. Akter, and M. A. Alam, "An ensemble machine learning approach to identify risk factors of anemia among infants and young children in developing nations," *Scientific Reports*, vol. 14, Art.no.1045, 2024, doi:10.1038/s41598-024-51290-y. Available:

<https://www.nature.com/articles/s41598-024-51290-y>

[22] R. Tibshirani, "Regression shrinkage and selection via the lasso," *Journal of the Royal Statistical Society: Series B (Methodological)*, vol. 58, no. 1, pp. 267–288, 1996, doi: 10.1111/j.2517-6161.1996.tb02080.x. Available: <https://rss.onlinelibrary.wiley.com/doi/abs/10.1111/j.2517-6161.1996.tb02080.x>

[23] P. Geurts, D. Ernst, and L. Wehenkel, "Extremely randomized trees," *Machine Learning*, vol. 63, no.1, pp.3–42, 2006, doi:10.1007/s10994-006-6226-1. Available:

<https://link.springer.com/article/10.1007/s10994-006-6226-1>

[24] H. T. Nguyen and A. M. Nguyen, "Applying gradient boosted decision trees to population-based healthcare indices," *Health Information Science and Systems*, vol. 11, Art. no. 18, 2023, doi: 10.1007/s13755-023-00224-x. Available: <https://link.springer.com/article/10.1007/s13755-023-00224-x>

[25] K. B. Shrestha, "Maternal and child risk factors associated with childhood anemia in South-East Asia: A predictive modeling study using demographic data," *Journal of Tropical Pediatrics*, vol. 70, no.3, Art.no.fmae015, 2024, doi:10.1093/tropej/fmae015. Available:

<https://academic.oup.com/tropej/article/70/3/fmae015/7665719>

[26] F. Pedregosa, G. Varoquaux, A. Gramfort, v. Michel, et al., "Scikit-learn: Machine learning in Python," *Journal of Machine Learning Research*, vol. 12, pp. 2825–2830, 2011. Available: <https://jmlr.org/papers/v12/pedregosa11a.html>

[27] C. Cortes and V. Vapnik, "Support-vector networks," *Machine Learning*, vol. 20, no. 3, pp. 273–297, 1995, doi:10.1007/BF00994018. Available: <https://link.springer.com/article/10.1007/BF00994018>

[28] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*, 2nd ed. New York, NY: Springer, 2009, doi: 10.1007/978-0-387-84858-7. Available: <https://link.springer.com/book/10.1007/978-0-387-84858-7>

[29] N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, "SMOTE: Synthetic minority over-sampling technique," *Journal of Artificial Intelligence Research*, vol. 16, pp. 321–357, 2002, doi: 10.1613/jair.953. Available: <https://www.jair.org/index.php/jair/article/view/10308>

[30] Y. Freund and R. E. Schapire, "A decision-theoretic generalization of on-line learning and an application to boosting," *Journal of Computer and System Sciences*, vol. 55, no. 1, pp. 119–139, 1997, doi:10.1006/jcss.1997.1504. Available:

<https://www.sciencedirect.com/science/article/pii/S002200009791504X>