

Application Of XGBoost-Based Machine Learning Methods To Predict Stunting

Muhammad Fariz Anhar ^{a,1,*}, Benfano Soewito ^{b,2}

^a Master of Computer Science, Bina Nusantara University, West Jakarta, Jakarta 11530, Indonesia

^b Master of Computer Science, Bina Nusantara University, West Jakarta, Jakarta 11530, Indonesia

¹ mfariz.anhar@gmail.com *; ² bsoewito@binus.edu

* corresponding author

ARTICLE INFO

Keywords

Machine Learning

Prediction

SMOTE

Stunting

XGBoost

ABSTRACT

Child stunting remains a major public-health challenge across Asia, impairing growth, cognition, and lifelong productivity. Early risk identification is critical, yet conventional screening offers limited predictive power and scalability. This study evaluates machine-learning approaches for stunting prediction using routinely collected infant data, proposing XGBoost and benchmarking it against Logistic Regression and Random Forest. An Asian infant dataset was compiled, label encoding and standardization were applied, class imbalance was addressed with SMOTE, the three models were trained and hyperparameter tuning was performed within a reproducible pipeline. Performance was assessed using Area Under the ROC Curve (AUC) and confusion matrices. XGBoost with SMOTE achieved the highest AUC (0.85), exceeding Random Forest (0.83) and Logistic Regression (0.73). Confusion-matrix analysis indicates that XGBoost separates stunted from non-stunted cases more effectively. Models trained without SMOTE performed worse, underscoring the value of imbalance correction. These findings suggest that ML assisted screening can enable earlier, data-driven risk stratification and targeted interventions. Practical deployment, however, may be constrained by the need for a GPU enabled computer and an IDE based workflow, motivating external validation and implementation refinement.

This is an open access article under the [CC-BY-SA](#) license.



1. Introduction

Stunting is a condition of impaired growth and development in children resulting from long-term nutritional deficiencies. According to the World Health Organization (WHO), stunting is a developmental disorder in children caused by poor nutrition, repeated infections, or inadequate psychosocial stimulation. A child is defined as stunted if their height for age is more than two standard deviations below the WHO Child Growth Standards.

Without the assistance of AI, detecting stunting may rely heavily on traditional, manual methods that can be time-consuming and may lead to delayed interventions. This can result in missed opportunities for early treatment, which is crucial for preventing long-term developmental issues.

With advancing technology, nearly all diseases can now be detected with the assistance of media such as computers.

Machine learning is a branch of artificial intelligence (AI) that focuses on developing systems capable of learning from data. Machine learning means a type of learning that greatly helps in solving problems, making it easier to accomplish tasks. In the healthcare field, machine learning makes it easier to accomplish tasks, for example, doctors can diagnose heart disease quickly without taking a long time. XGBoost is an improvement on the Gradient Boosting algorithm. XGBoost was introduced by Chen and Guestrin in 2016.

XGBoost, short for eXtreme Gradient Boosting, is an ensemble learning algorithm that uses boosting techniques, introduced by Tianqi Chen in 2014. This algorithm focuses more intensively on samples that were previously misclassified, following the basic concept of gradient boosting. In XGBoost, there are two types of trees used, regression trees and classification trees.

SMOTE stands for Synthetic Minority Over-sampling Technique. It's a popular method used in machine learning to address the issue of imbalanced datasets. When a dataset has a significant imbalance between the majority and minority classes, traditional machine learning algorithms tend to be biased towards the majority class.

SMOTE works by creating synthetic samples of the minority class, thereby balancing the dataset. It does this by selecting minority class instances and generating new synthetic data points along the line segments joining the k-nearest neighbors of each minority class instance. This helps to improve the performance of machine learning models on imbalanced datasets by giving the minority class a more equal representation. Research is needed to produce a Machine Learning model capable of predicting stunting in infants quickly and using modern methods.

The reviewed literature consistently positions XGBoost as a high-performing baseline for tabular prediction across domains, achieving low error in traffic accident forecasting, strong discrimination in clinical survival prediction when rigorously tuned, and state-of-the-art performance on imbalanced datasets when paired with resampling techniques such as SMOTE.

Studies on heart failure show that hyperparameter optimization materially elevates effectiveness with tree-structured Parzen estimators and Bayesian search often outperforming grid or random search while work on liver disease demonstrates that combining SMOTE with Bayesian tuning yields the best AUC among competing configurations. Comparative analyses in hepatitis C stratification further indicate that model choice interacts with phenotype, where SMOTE–Random Forest favored fibrosis whereas SMOTE–XGBoost excelled for cirrhosis, underscoring the need to match algorithms to clinical subproblems.

Beyond healthcare, SMOTE-augmented XGBoost attains very high accuracy, precision, recall, and ROC-AUC for air-quality classification, reinforcing its robustness on imbalanced data. Taken together, these findings motivate three design principles for this study: mitigate class imbalance perform systematic hyperparameter tuning, and benchmark against strong tree-based ensembles while accounting for task-specific error trade-off, principles that justify adopting and carefully optimizing XGBoost in the present work.

2. Method

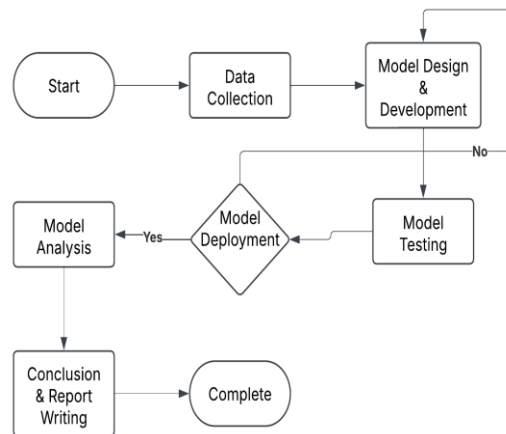


Fig. 1.Research Stages

2.1. Data Collection

The data collection stage formed the foundation of this research. For this study, the dataset was obtained from the Kaggle platform, under the title *Stunting_dataset*. The dataset was uploaded by a Kaggle contributor named Muhtarom, who designed it as an openly accessible resource for supervised machine learning tasks. The dataset is tabular in structure and includes information relevant to assessing child growth and nutritional health. Specifically, it contains attributes such as Gender, Age, Birth Weight, Birth Length, Body Weight, Body Length, and Breastfeeding status, with the target label being whether the child is stunted or not. In total, the dataset includes 10,000 rows and 8 columns, making it sufficiently large for machine learning experiments.

The file format is CSV (Comma Separated Values), which facilitates preprocessing and ensures compatibility with Python-based data science libraries such as pandas and NumPy. Each row in the file represents an individual child, and each column corresponds to one of the recorded attributes. Initial exploration of the dataset revealed certain issues, including imbalances between stunted and non-stunted cases. Approximately 7,955 rows indicated children who were stunted, while only 2,045 rows represented non-stunted children. This significant imbalance necessitated the use of balancing techniques such as SMOTE to avoid bias in training.

The dataset, while comprehensive in anthropometric data, did not provide socioeconomic or environmental features. Although this posed a limitation, it simplified the initial focus on biological and nutritional markers of stunting. The dataset was also clean, containing minimal missing values. For the purposes of this study, exploratory data analysis (EDA) was performed to examine variable distributions, check correlations among features, and detect anomalies or outliers. These steps were essential to ensure the integrity of the data before advancing to model development.

2.2. Model Design & Development

The second stage involved the careful design and development of predictive models. The XGBoost algorithm was selected as the primary classifier because of its efficiency and suitability for tabular structured data. XGBoost is an advanced implementation of gradient boosting decision trees that incorporates regularization terms to prevent overfitting. In this study, it was expected to outperform simpler linear models such as Logistic Regression and ensemble methods such as Random Forest.

The development process was implemented in Python using Google Colab as the execution environment. Google Colab provides access to high-performance computing resources, including GPU acceleration, which is beneficial for training complex models on large datasets. The workflow was structured as a pipeline consisting of four major components: preprocessing, resampling, training, and evaluation.

One of the critical challenges in the dataset was class imbalance. The Synthetic Minority Oversampling Technique (SMOTE) was therefore incorporated into the workflow. SMOTE operates by selecting minority class samples (in this case, non-stunted children) and generating new synthetic

instances along the vector lines connecting these points with their k-nearest neighbors. This approach increases the representation of the minority class, reducing the bias of classifiers toward the majority class. Additionally, Tomek Links were used in combination with SMOTE, resulting in the SMOTET technique. Tomek Links identify borderline samples in the majority class and remove them, thereby cleaning the decision boundaries. Together, SMOTET ensures that models are trained on balanced and representative data.

After preprocessing and balancing, the models were designed and implemented. XGBoost was fine-tuned by experimenting with different hyperparameters such as learning rate, number of trees (n_estimators), maximum tree depth (max_depth), subsample ratios, and minimum child weight. Logistic Regression was included as a baseline model, providing interpretability and simplicity but limited in its ability to capture complex relationships. Random Forest, another ensemble tree-based algorithm, was implemented to compare against XGBoost, particularly in its handling of nonlinearities and variable importance measures.

2.3. Model Testing

The testing stage was designed to rigorously assess the models under standardized conditions. The dataset was divided into training and testing subsets in an 80:20 ratio. This ensured that the majority of the data contributed to model learning while reserving a significant portion for unbiased evaluation. The testing phase focused on determining whether the models could generalize effectively to unseen data.

Success in testing was measured by a combination of metrics. Accuracy, while straightforward, was not considered sufficient on its own due to the dataset imbalance. Therefore, additional metrics such as Precision, Recall, and F1-Score were emphasized. Precision measures the proportion of true positives among all predicted positives, which is crucial in avoiding false alarms when predicting stunting. Recall measures the proportion of actual positives correctly predicted, which is critical in healthcare to avoid missing stunted cases. F1-Score, as the harmonic mean of Precision and Recall, balances the two concerns. A Confusion Matrix was generated to provide a detailed breakdown of prediction outcomes. Finally, the ROC-AUC metric was applied to evaluate the overall discriminatory power of each model.

To ensure robustness, cross-validation was also employed during training. This involved partitioning the training dataset into multiple folds and training the model across these folds to reduce variance in performance estimates. Hyperparameter tuning was integrated with cross-validation, ensuring that parameter configurations leading to overfitting were avoided.

2.4. Model Analysis

The model analysis phase focused on interpreting and improving the results. Each model was assessed using its Confusion Matrix and ROC curves. For example, XGBoost's confusion matrix revealed strong predictive ability, with high true positive and true negative values. Logistic Regression, in contrast, demonstrated limitations by misclassifying a higher number of cases due to its linear assumption. Random Forest performed relatively well but was less efficient and slightly less accurate than XGBoost.

The ROC curve for XGBoost showed an area under the curve (AUC) of 0.85, highlighting its superior ability to distinguish between stunted and non-stunted infants. Random Forest followed closely with an AUC of 0.83, while Logistic Regression lagged.

3. Results and Discussion

3.1. Data Collection

The dataset used in this study, titled "Stunting_dataset," is a tabular supervised-learning corpus comprising 10,000 rows and 8 columns with clinically relevant attributes: Gender, Age, Birth Weight, Birth Length, Body Weight, Body Length, and Breastfeeding.

These features capture perinatal and anthropometric factors commonly associated with childhood growth outcomes and thus provide a practical basis for stunting risk stratification at scale. An illustrative preview of the first 10 rows is presented to contextualize the schema and value ranges that informed preprocessing and modeling decisions.

	A	B	C	D	E	F	G	H
1	Gender, Age, Birth Weight, Birth Length, Body Weight, Body Length, Breastfeeding, Stunting							
2	Male, 17, 3, 49, 10, 72.2, No, No							
3	Female, 11, 2, 9, 49, 2, 9, 65, No, Yes							
4	Male, 16, 2, 9, 49, 8, 5, 72.2, No, Yes							
5	Male, 31, 2, 8, 49, 6, 4, 63, No, Yes							
6	Male, 15, 3, 1, 49, 10, 5, 49, No, Yes							
7	Female, 11, 2, 8, 49, 8, 5, 65, No, No							
8	Male, 35, 2, 8, 49, 10, 5, 72.2, No, Yes							
9	Female, 17, 2, 8, 49, 8, 63, No, Yes							
10	Female, 10, 2, 7, 49, 8, 4, 73.5, No, No							
11	Female, 16, 2, 8, 49, 8, 5, 65, No, Yes							

Fig. 2. First 10 Rows of the Dataset

3.2. Model Design

The modeling pipeline was implemented in Python within the Google Colab environment and followed a reproducible sequence: library and dataset setup, data transformation and preprocessing, feature-label separation, stratified train-test split (80/20), class-imbalance handling, model training, and out-of-sample evaluation.

Categorical variables were label-encoded, and features were standardized to harmonize scale across variables prior to learning, ensuring consistent inputs for both linear and tree-based estimators under a unified pipeline. To address class imbalance, Synthetic Minority Over-sampling Technique (SMOTE) further paired with Tomek links (SMOTET) was applied to the training set to synthesize minority-class samples and clean borderline pairs without contaminating the test distribution.

Hyperparameter tuning was then performed to explore model configuration spaces, followed by the construction of three comparative classifiers as Logistic Regression, Random Forest, and the gradient-boosted ensemble XGBoost and the preparation of an inference function to support predictions on new data instances.

3.3. Model Analysis

This stage involves analyzing the models that have been created. There are 3 models using 3 different algorithms. The first prediction model uses XGBoost, the second uses Logistic Regression, and the last uses Random Forest. The results of these three models with different algorithms are shown.

Table 1.		Classification Report		
Algorithm	Precision	Recall	F1 Score	Accuracy
XGBoost	0.84	0.85	0.84	0.85
Logistic Regression	0.80	0.73	0.75	0.73
Random Forest	0.82	0.83	0.82	0.83

The table indicates that XGBoost attains the strongest overall performance on this dataset, with Accuracy 0.85, weighted Precision 0.84, weighted Recall 0.85, and weighted F1-score 0.84, while Random Forest is close behind (Accuracy 0.83) and Logistic Regression is substantially lower (Accuracy 0.73).

Accuracy quantifies the proportion of all predictions that are correct and is computed as in

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

serving as a coarse indicator of model quality especially when classes are balanced.

Precision measures the reliability of positive predictions, as in

$$Precision = \frac{TP}{TP+FP} \quad (2)$$

emphasizing the control of false positives.

Recall measures the ability to capture actual positives, as in

$$Recall = \frac{TP}{TP+FN} \quad (3)$$

emphasizing the control of false negatives.

The F1-score summarizes the balance between precision and recall via the harmonic mean, as in

$$F1 = \frac{2 \times Precision \times Recall}{Precision+Recall} \quad (4)$$

and is particularly informative when classes are imbalanced or error costs differ.

To reflect class prevalence, the table reports weighted averaging, which aggregates per-class metrics using class support so that each class contributes proportionally to its frequency, as in

$$Weighted = \frac{\sum_i w_i x_i}{\sum_i w_i} \quad (5)$$

This approach avoids overstating performance on majority classes and provides a faithful overall summary under skewed distributions, though it can still mask weaknesses on very rare classes if those classes have minimal weight.

Interpreting the rows, XGBoost's weighted Precision 0.84, Recall 0.85, and F1 0.84 indicate a balanced trade-off, few false alarms and few misses, which explains the strong F1 mirroring the similarly high Precision and Recall.

Random Forest shows slightly reduced Precision 0.82 and Recall 0.83 relative to XGBoost, producing a marginally lower F1 of 0.82 and Accuracy of 0.83, consistent with a small increase in either false positives or false negatives. Logistic Regression exhibits the largest deficit in Recall at 0.73, which depresses its F1 to 0.75 and Accuracy to 0.73, indicating a greater tendency to miss true positives compared with the tree-based methods.

The performance gaps are practically meaningful, XGBoost exceeds Logistic Regression by 0.12 in Accuracy and by 0.09 in F1, and it edges Random Forest by 0.02 in both Accuracy and Recall, highlighting a consistent advantage across correctness and sensitivity. These differences imply that model selection should favor the method that keeps Precision and Recall simultaneously high, because F1 will track the worse of the two and penalize imbalanced error profiles.

From an algorithmic standpoint, gradient-boosted trees as implemented in XGBoost incorporate regularization and pruning in the objective to improve generalization, and they capture complex, non-linear interactions common in structured/tabular data, which often leads to superior accuracy over linear baselines. This property offers a plausible explanation for the observed ranking in the table, where the boosted ensemble outperforms the linear Logistic Regression and slightly improves upon the bagged Random Forest.

For transparent reporting, it is advisable to present Accuracy alongside Precision, Recall, and F1 and to interpret them jointly, since metric choice hinges on class balance and the relative costs of false positives versus false negatives. When class imbalance is present, continue to report weighted averages for an overall picture and consider supplementing with macro-averaged or per-class scores

in an appendix to ensure minority-class performance is visible to readers. Furthermore, there is a visualization of the Receiver Operating Characteristic (ROC) from the code.

The ROC-Curve shows the trade-off between the True Positive Rate (TPR) (sensitivity) and the False Positive Rate (FPR) for various thresholds. A better model will have a curve closer to the upper left corner, indicating an ability to effectively separate the classes. This ROC-curve is shown in the image below.

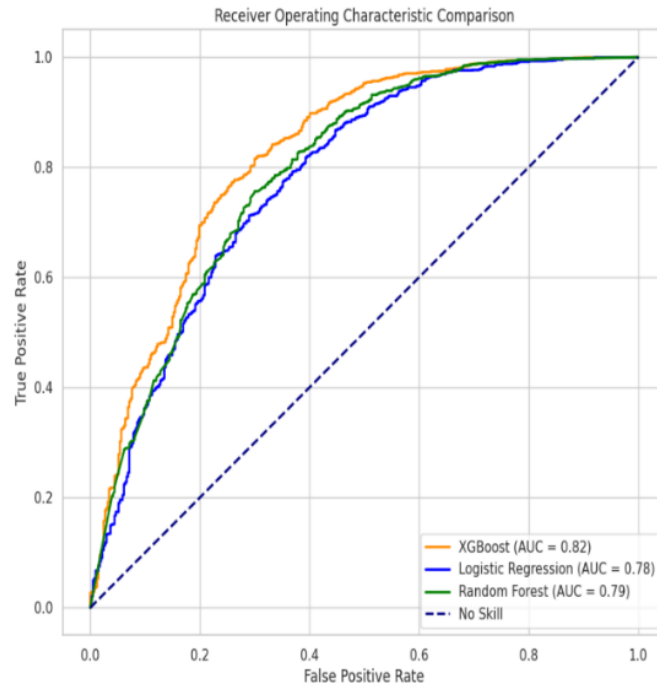


Fig. 3. ROC-Curve Model

Based on this ROC-Curve, XGBoost is the best model according to the ROC-Curve, with the highest AUC (0.82), indicating it can distinguish between classes very well.

Random Forest also performs well but slightly falls short of XGBoost. Logistic Regression has the lowest performance but still far exceeds the baseline. Receiver Operating Characteristic (ROC) analysis corroborated these findings, with XGBoost exhibiting the most favorable trace and the largest Area Under the Curve (AUC = 0.82), indicating strong class separability across decision thresholds relative to competing baselines.

This ranking advantage is consistent with the stage-wise error correction and regularization properties of gradient boosting, which together enhance discrimination under feature heterogeneity and residual class imbalance. This ROC Curve supports the results from the classification report, indicating that XGBoost is the best choice for use in this model prediction.

Additionally, there is a Confusion Matrix. The Confusion Matrix is a table used to evaluate the performance of a classification model. This table compares the model's predictions with the actual values in the dataset, providing details on the number of correct and incorrect predictions for each class. The Confusion Matrix for the XGBoost model is shown in the following image.



Fig. 4.Confusion Matrix for XGBoost

True Positive (TP) with a value of 214. This means that the model correctly predicted positive (class 1) when the actual data was also positive. Then, False Positive (FP) with a value of 193. This means the model predicted positive (class 1) but the actual data was negative (class 0).

This error is known as a false alarm. Next, False Negative (FN) with a value of 104. This means the model predicted negative (class 0) but the actual data was positive (class 1). This error is known as a missed detection. Lastly, True Negative (TN) with a value of 1489. This means the model correctly predicted negative (class 0) when the actual data was also negative. These counts reflect a model that distinguishes positive and negative classes reliably while revealing actionable trade-offs, opportunities to reduce missed detections (FN) through threshold calibration, cost sensitive learning, or targeted feature refinements in deployment contexts where recall is paramount.

Taken together, the quantitative results and diagnostic plots substantiate XGBoost as the preferred classifier for this application, offering the best composite of accuracy, ranking quality, and error balance after SMOTE-based resampling under the same features and protocol. The confusion matrix shows that XGBoost performs well in distinguishing between positive and negative classes. Therefore, this study uses the model from the XGBoost algorithm. The model achieved an overall accuracy of 85%, indicating that it correctly predicted whether a child was stunted or not in 85% of cases.

4. Conclusion

This study designed, implemented, and evaluated a supervised machine learning pipeline for stunting prediction centered on an XGBoost classifier and demonstrated that tabular, routinely collected attributes Gender, Age, Birth Weight, Birth Length, Body Weight, Body Length, and Breastfeeding, are sufficient to yield strong predictive performance under a principled preprocessing and evaluation protocol. The end-to-end process included categorical encoding, feature scaling, an 80/20 train-test split, explicit class-imbalance handling via SMOTE (and Tomek links where applicable), comparative baselines, and out-of-sample assessment, thereby enabling a fair comparison across models and a transparent accounting of the trade-offs relevant to early risk screening in pediatric contexts.

Empirically, the XGBoost model outperformed the Logistic Regression and Random Forest baselines on the same features and processing steps, achieving an overall accuracy of 85% and a ROC-AUC in the low-to-mid 0.80s (0.82–0.85 depending on evaluation), which jointly indicate good ranking quality and reliable class separability on a dataset exhibiting real-world class skew.

Complementary diagnostics using the confusion matrix (TP = 214, FP = 193, FN = 104, TN = 1,489 on the held-out test set) clarified operational behavior at a nominal threshold, highlighting both effective overall discrimination and specific error modes—particularly missed stunted cases—that matter for programmatic decision-making and justify calibrated operating points when sensitivity is prioritized. Taken together, the findings substantiate the choice of a gradient-boosted ensemble as the primary modeling approach for this problem setting, with stage-wise error correction, embedded regularization, and flexible tree-based partitions proving advantageous over linear decision boundaries and bagging-only ensembles on heterogeneous clinical tabular data.

At the same time, several limitations temper direct generalization and point to concrete avenues for improvement in future iterations of the system. First, the presence of class imbalance, even after resampling, can still interact with default thresholds to bias decisions toward the majority class, making false negatives for the stunted class a salient risk in screening workflows unless thresholding and loss weighting are explicitly tuned to program goals. Second, reliance on a single dataset and data-generation process constrains external validity, motivating independent validation across cohorts, geographies, and measurement practices, as well as robustness checks under shifts in prevalence and socio-demographic distributions to ensure transportability and equity. Third, while inference can run on modest hardware, practical retraining, large-scale evaluation, and iterative experimentation benefit from access to accelerated compute and a consistent software stack, and the current research pipeline still presumes a development environment before streamlined deployment to non-specialist users.

Guided by these observations and the model diagnostics, targeted enhancements are feasible on the model, data, and evaluation fronts without disrupting the core pipeline. On the model side, class-weighted or cost-sensitive learning, threshold calibration based on validation data, and focused hyperparameter exploration, particularly around learning rate, maximum depth, minimum child weight, and subsampling ratios—are well-positioned to improve the bias variance balance and increase recall for the stunted class while controlling false alarms. On the data side, curating complementary features alongside refined anthropometrics, standardizing preprocessing, and carefully handling missingness can strengthen signal and reduce measurement noise, provided additions remain feasible within routine data collection practices. On the evaluation side, consistently reporting per-class metrics, probability calibration quality, and confusion matrices under multiple thresholds will make trade-offs explicit and help ensure that gains in sensitivity do not unduly degrade calibration or inflate false positives in resource-constrained settings.

From an application perspective, the current results support the use of the XGBoost-centered pipeline as a practical decision support component for early stunting risk stratification, subject to deployment policies that align operating thresholds with program objectives and to safeguards that preserve data privacy, security, and accountability in real-world use. Packaging the trained model behind an accessible interface, such as a web application or lightweight client, using exactly the same preprocessing steps as in training would extend usability beyond specialist environments and enable batch or single-record prediction, while retaining fidelity and interpretability for clinicians and public health practitioners. Importantly, any production roll-out should incorporate model versioning, scheduled re-evaluations, and monitoring for distribution shifts to sustain accuracy and fairness over time, and should be accompanied by clear documentation on how to tune thresholds for recall or precision depending on screening priorities and operational constraints.

In summary, this research demonstrates that an XGBoost-based approach, coupled with disciplined preprocessing and class-imbalance handling, provides a robust and effective foundation for stunting prediction on tabular health data, outperforming common alternatives under identical conditions and yielding accuracy around 85% with ROC-AUC in the low-to-mid 0.80s on a held-out test set. Future work that expands feature breadth in a controlled and interpretable manner, validates across independent cohorts, refines cost-sensitive objectives and threshold policies, and operationalizes delivery through accessible software can materially enhance accuracy, generalizability, and impact, positioning the system for reliable, equitable, and scalable support of early stunting detection and targeted intervention pathways.

Acknowledgment

The author expresses sincere gratitude to Benfano Soewito, the second author and thesis supervisor, Master of Computer Science, Bina Nusantara University, Jakarta, Indonesia, for continuous guidance and constructive feedback that enabled the completion of this article.

Declarations (HEADING 5)

Author contribution. [First Author;Muhammad Fariz Anhar] Conceptualization, Methodology, Writing–original draft, Editing, Supervision. [Second Author;Benfano Soewito] Investigation. Data curation. Writing–original draft, Review, Editing.

Funding statement. This research received no specific grant from any funding agency in the public, commercial, or not for profit sectors.

Conflict of interest. The authors declare that they have no conflict of interest.

Additional information. No additional information is available for this paper.

References

- [1] Ali, S., Khorrami, B., Jehanzaib, M., Tariq, A., Ajmal, M., Arshad, A., Shafeeque, M., Dilawar, A., Basit, I., Zhang, L., Sadri, S., Niaz, M. A., Jamil, A., & Khan, S. N. (2023). Spatial Downscaling of GRACE Data Based on XGBoost Model for Improved Understanding of Hydrological Droughts in the Indus Basin Irrigation System (IBIS). *Remote Sensing*, 15(4). <https://doi.org/10.3390/rs15040873>
- [2] Andriansyah, D.-, & Eka Wulansari Fridayanthie. (2023). Optimization of Support Vector Machine and XGBoost Methods Using Feature Selection to Improve Classification Performance. *JOURNAL OF INFORMATICS AND TELECOMMUNICATION ENGINEERING*, 6(2), 484–493. <https://doi.org/10.31289/jite.v6i2.8373>
- [3] Aurima, J., Susaldi, S., Agustina, N., Masturoh, A., Rahmawati, R., & Tresiana Monika Madhe, M. (2021). Faktor-Faktor yang Berhubungan dengan Kejadian Stunting pada Balita di Indonesia. *Open Access Jakarta Journal of Health Sciences*, 1(2), 43–48. <https://doi.org/10.53801/oajjhs.v1i3.23>
- [4] Chen, L., Chen, P., & Lin, Z. (2020). Artificial Intelligence in Education: A Review. *IEEE Access*, 8, 75264–75278. <https://doi.org/10.1109/ACCESS.2020.2988510>
- [5] Fitri, L. (2018). HUBUNGAN BBLR DAN ASI EKSLUSIF DENGAN KEJADIAN STUNTING DI PUSKESMAS LIMA PULUH PEKANBARU. *Jurnal Endurance*, 3(1), 131. <https://doi.org/10.22216/jen.v3i1.1767>
- [6] Hajek, P., Abedin, M. Z., & Sivarajah, U. (2023). Fraud Detection in Mobile Payment Systems using an XGBoost-based Framework. *Information Systems Frontiers*, 25(5), 1985–2003. <https://doi.org/10.1007/s10796-022-10346-6>
- [7] Juhdan Abdullah M, F. A. J. M. I. A. D. I. S. (2023). Hubungan Perkembangan Teknologi AI Terhadap Pembelajaran Mahasiswa. *Jurnal Pendidikan Seroja*, 4(2).
- [8] Komalasari, E. S. R. S. H. I. (2020). Faktor-faktor Penyebab Kejadian Stunting Pada Balita. *Majalah Kesehatan Indonesia*, 1(2).
- [9] Lestari, I., Akbar, M., & Intan, B. (2023). Perbandingan Algoritma Machine Learning Untuk klasifikasi Amenorrhea. *Journal of Computer and Information Systems Ampera* 4(1).
- [10] Marlin, K., Faisal, F.R., Noviandy. (2023). Manfaat dan Tantangan Penggunaan Artificial Intelligences (AI) Chat GPT Terhadap Proses Pendidikan Etika dan Kompetensi Mahasiswa di Perguruan Tinggi. *Journal of Science Research*, 3(6).
- [11] Maulana, A., Faisal, F. R., Noviandy, T. R., Rizkia, T., Idroes, G. M., Tallei, T. E., El-Shazly, M., & Idroes, R. (2023). Machine Learning Approach for Diabetes Detection Using Fine-Tuned XGBoost Algorithm. *Infolitika Journal of Data Science*, 1(1), 1–7.
- [12] Nababan, A. A., Jannah, M., Aulina, M., & Andrian, D. (2023). PREDIKSI KUALITAS UDARA MENGGUNAKAN XGBOOST DENGAN SYNTHETIC MINORITY OVERSAMPLING TECHNIQUE (SMOTE) BERDASARKAN INDEKS STANDAR PENCEMARAN UDARA (ISPU). *Jurnal Teknik Informatika Kaputama (JTik)*, 7(1).

-
- [13] Noviandy, T. R., Idroes, G. M., Maulana, A., Hardi, I., Ringga, E. S., & Idroes, R. (2023). Credit Card Fraud Detection for Contemporary Financial Management Using XGBoost-Driven Machine Learning and Data Augmentation Techniques. *Indatu Journal of Management and Accounting*, 1(1), 29–35.
- [14] Rizky Mubarak, M., Herteno, R., Komputer Fakultas Matematika dan Ilmu Pengetahuan Alam Universitas Lambung Mangkurat Jalan Ahmad Yani Km, I., & Selatan, K. (n.d.). HYPER-PARAMETER TUNING PADA XGBOOST UNTUK PREDIKSI KEBERLANGSUNGAN HIDUP PASIEN GAGAL JANTUNG.
- [15] Xu, K., Han, Z., Xu, H., & Bin, L. (2023). Rapid Prediction Model for Urban Floods Based on a Light Gradient Boosting Machine Approach and Hydrological–Hydraulic Model. *International Journal of Disaster Risk Science*, 14(1), 79–97.