

Article

Not peer-reviewed version

Enhancing the Interpretability of Malaria and Typhoid Diagnosis with Explainable AI and Large Language Models

[Kingsley Attai](#)*, [Moses Ekpenyong](#), [Constance Amannah](#), [Daniel Asuguo](#), Peterben Ajuga, [Okure Obot](#), Ekemini Johnson, Anietie John, Omosivie Maduka, [Christie Akwaowo](#), [Faith-Michael Uzoka](#)

Posted Date: 14 August 2024

doi: 10.20944/preprints202408.0966.v1

Keywords: Malaria Diagnosis; Typhoid Diagnosis; Machine Learning; XAI; LIME; GPT; BERT; ChatGPT; Gemini; Perplexity; Explainability; Interpretability



Preprints.org is a free multidiscipline platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This is an open access article distributed under the Creative Commons Attribution License which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Article

Enhancing the Interpretability of Malaria and Typhoid Diagnosis with Explainable AI and Large Language Models

Kingsley Attai ^{1,*}, Moses Ekpenyong ², Constance Amannah ³, Daniel Asuquo ², Peterben Ajuga ⁴, Okure Obot ², Ekemini Johnson ¹, Anietie John ¹, Omosivie Maduka ⁵, Christie Akwaowo ⁶ and Faith-Michael Uzoka ⁷

¹ Department of Mathematics and Computer Science, Ritman University, Ikot Ekpene 530101, Nigeria; eke5461@gmail.com (E.J.); aniettejhn5@gmail.com (A.J.)
² Department of Computer Science, Faculty of Science, University of Uyo, Uyo 520103, Nigeria; mosesekpenyong@uniuyo.edu.ng (M.E); danielasuquo@uniuyo.edu.ng (D.A); okureobot@uniuyo.edu.ng (O.O)
³ Department of Computer Science, Ignatius Ajuru University of Education, Port Harcourt 500102, Nigeria; aftermymisc@gmail.com
⁴ Department of Computer Engineering, Faculty of Engineering, Gregory University, Uturu 441106, Nigeria; ajugapeterben@gmail.com
⁵ University of Port Harcourt Teaching Hospital, Port Harcourt 500102, Nigeria; omosivie.maduka@gmail.com
⁶ University of Uyo Teaching Hospital, Uyo, Uyo 520103, Nigeria; christieakwaowo@uniuyo.edu.ng
⁷ Department of Mathematics and Computing, Mount Royal University, Calgary, AB T3E 6K6, Canada; fuzoka@mtroyal.ca
* Correspondence: attai.kingsley@ritmanuniversity.edu.ng; Tel.: +2348101250218

Abstract: Malaria and Typhoid fever are prevalent diseases in tropical regions, exacerbated by unclear protocols, drug resistance, and environmental factors. Prompt and accurate diagnosis is crucial to improve accessibility and reduce mortality rates. Traditional diagnosis methods cannot effectively capture the complexities of these diseases due to the presence of confusable symptoms. Although machine learning (ML) models offer accurate predictions, they operate as "black boxes" with non-interpretable decision-making processes, making it challenging for healthcare providers to comprehend how the conclusions are reached. This study employs explainable AI (XAI) models such as Local Interpretable Model-agnostic Explanations (LIME), and Large Language Models (LLMs) like ChatGPT, Gemini, and Perplexity, to clarify diagnostic results for healthcare workers, building trust and transparency in medical diagnostics. The models were implemented on Google Colab and Visual Studio Code because of their rich libraries and extensions. Results showed that the Random Forest model outperformed Extreme Gradient Boost and Support Vector Machines and important features for predicting Malaria and Typhoid were identified with the LIME plots. Among LLMs, ChatGPT 3.5 demonstrates a comparative advantage over Gemini and Perplexity, hence highlighting how crucial it is to use a hybrid strategy that combines ML, XAI, and LLMs to improve diagnostic performance and reliability in healthcare applications. The study recommends implementing the integrated ML, LIME, and LLM processes because it streamlines the development and maintenance process's workflow overall, makes use of resources more effectively, and improves explainability by allowing the LLM to consider the complete context that LIME provides.

Keywords: malaria diagnosis; typhoid diagnosis; machine learning; XAI; LIME; GPT; BERT; ChatGPT; gemini; perplexity; explainability; interpretability

1. Introduction

Typhoid fever and malaria are two of the most prevalent febrile diseases in the world, presenting serious public health issues, especially in tropical and subtropical areas. Typhoid and malaria are common in these areas due to the high humidity, temperatures, inadequate healthcare facilities, and the shortage of qualified healthcare providers [1]. *Salmonella enterica* serotype Typhi is the bacteria that causes typhoid fever or enteric fever which affects millions of people worldwide and can have

serious consequences if left untreated [2–4]. Malaria on the other hand is caused by plasmodium parasites that are transmitted by Anopheles mosquito bites, infecting millions of people and claiming the lives of hundreds of thousands every year [5–7]. Malaria is the second most studied disease [8] due to its widespread prevalence, high mortality rate, drug resistance, and environmental factors such as climate change. The prompt and effective diagnosis of these febrile diseases is essential for efficient treatment and care, but current diagnostic techniques often face limitations in accessibility, specificity, and sensitivity.

Machine learning (ML) algorithms are frequently used in the healthcare sector, to help decision-makers make well-informed decisions [9,10]. Medical diagnostics has found ML machine learning to be a potent tool that can improve the efficiency and accuracy of diagnosis, but to guarantee that medical professionals can rely on and comprehend the judgments made by these models, the use of ML models in clinical settings demands a high level of interpretability and transparency. According to Anderson & Thomas [11], concerns about ML algorithms' lack of interpretability frequently impede their acceptance in the healthcare sector. Their capacity to comprehend and interpret the choices made by ML models is critical in this sector, as decisions can have a significant impact on patient outcomes. To address this challenge, an explainable AI (XAI) technique like Local Interpretable Model-agnostic Explanations (LIME) offers insights into how models arrive at their predictions, thereby promoting trust and aiding in clinical decision-making by healthcare professionals. XAI is becoming increasingly important in the healthcare sector, where making decisions has extremely high stakes because it is challenging for healthcare professionals to trust and comprehend the decisions made by traditional machine learning models. In clinical settings, where comprehension of the reasoning behind a diagnosis is critical for patient safety, regulatory compliance, and ethical considerations, the lack of interpretability may impede the adoption of AI [12]. Therefore, XAI offers solutions to these problems by facilitating AI models' transparent and intelligible decision-making process. XAI techniques such as LIME are widely utilized to clarify the inner workings of complex models. LIME operates by using an interpretable model local to the prediction to approximate the black-box model. It modifies the input data, tracks how the predictions change as a result, and then fits a straightforward, understandable model to these modified samples [13,14]. In situations where individual case explanations are required, LIME is especially helpful as it helps determine which characteristics are most important for a particular prediction. The interpretability of ML models in the healthcare industry is greatly enhanced by LIME, which allows physicians to better comprehend and rely on AI-driven insights and their capacity to offer concise, useful explanations improves the usefulness of AI systems in the processes of diagnosis and treatment planning. LIME has been applied in several healthcare settings such as in diagnosing diabetes [15], classification of co-morbidities associated with febrile diseases in children and pregnant women [16], and transparent health predictions [17]. To further improve accuracy and explainability, incorporating large language models (LLMs) into diagnostic processes seems promising in addition to XAI techniques. These models can bridge the gap between complex ML algorithms and clinical understanding. They are trained on a wealth of medical data and can provide distinctive interpretations and generate detailed, contextually relevant explanations for diagnostic outcomes.

The use of LLMs in medical contexts has advanced significantly; thanks to projects like Generative Pre-trained Transformer (GPT) and Bidirectional Encoder Representations from Transformers (BERT). These models can produce human-like text and comprehend intricate linguistic patterns because they have been trained on enormous volumes of text data. The applications of BERT go beyond identifying pandemic illnesses; it can also be used to process electronic medical records and evaluate the results of goals-of-care talks in clinical trials [18–20]. GPT has proven to be remarkably adept at producing coherent and contextually relevant text in various domains [21]. GPT can help in the healthcare industry by delivering comprehensive patient reports, producing justifications for medical diagnoses, and offering assistance during clinical decision-making processes [22]. The accuracy and explainability of diagnostic systems can be greatly improved by integrating these LLMs and they can produce thorough narratives that clarify the reasoning behind diagnostic predictions, which facilitates clinician comprehension and validation of

AI recommendations. This ability is essential for bridging the gap between cutting-edge AI models and real-world, routine clinical use, raising the standard of healthcare delivery as a whole.

This study aims to enhance the interpretability of typhoid and malaria diagnosis using ML techniques like Extreme gradient boost (XGBoost), Random Forest (RF), and Support Vector Machine (SVM) with LIME, and LLMs such as GPT, Gemini, and Perplexity. It emphasizes the potential of integrating these tools to interpret and contextualize medical data, hence, bridging the gap between healthcare worker comprehension and complex ML diagnoses. Real-time patient dataset consisting of symptoms and diagnoses of malaria and typhoid was collected from healthcare facilities across the Niger Delta region of Nigeria. By leveraging these advanced tools, we seek to develop a diagnostic model that delivers precise diagnoses and provides transparent and understandable insights into their decision-making processes. This research holds significant potential to improve diagnostic practices, ultimately contributing to better patient outcomes and advancing the field of medical diagnostics. This study can advance the field of diagnostic medicine and enhance diagnostic procedures, which will ultimately lead to better patient outcomes. This study's primary contributions are:

- The consideration of multiple diseases (typhoid fever and malaria) allows for a thorough evaluation of the patient's health, which is essential for managing co-infection and comorbidity.
- Using real-world data ensures that the models are trained and validated on clinical cases thereby enhancing the practical relevance and applicability of our findings.
- The black-box nature of many ML models is addressed by the integration of an XAI method, which gives medical professionals transparent and comprehensible insights into how each feature influences the diagnosis, ensuring that diagnostic results are presented in a way that is meaningful for easier interpretation. By focusing on interpretability, healthcare workers can make more accurate and timely diagnoses.
- LLMs give the diagnosis process an extra layer of context-aware understanding and incorporating them makes it possible to better understand and analyze complex medical outcomes.
- The combination of LLMs and conventional ML models enables a thorough comparison of various diagnosis strategies. This not only demonstrates the models' efficacy but also the advantages and disadvantages of each approach to managing medical data.
- The integration of XAI, LLMs, and ML puts this work at the forefront of medical AI research. It demonstrates the viability and benefits of using a hybrid approach to address difficult diagnostic problems, establishing a standard for further study in the area.

The study is prepared as follows: The methodology is presented in Section 2, including data collection, preprocessing, and the application of XAI and ML models, along with the incorporation of LLMs for improved diagnostic interpretability. The results are discussed in Section 3, evaluating the effectiveness of various algorithms and illustrating how XAI offers insights into model decisions, along with the implications for clinical practice. Section 4 concludes the study, highlighting its limitations, and offering recommendations for further research to advance diagnostic techniques.

2. Methodology

2.1. Malaria and typhoid diagnosis framework

The proposed diagnosis framework for malaria and typhoid fever is presented in Figure 1. The major components of the framework include a healthcare worker, medical experts, mobile device for the collection, processing, and storage of information locally and on a cloud-based storage for decision making. Patient data was obtained from medical experts and pre-processed into a format suitable for machine learning modeling and processing by large language models. The proposed model can be utilized in the diagnoses of typhoid fever and malaria with enhanced accuracy and explainability by a healthcare worker through a mobile device.

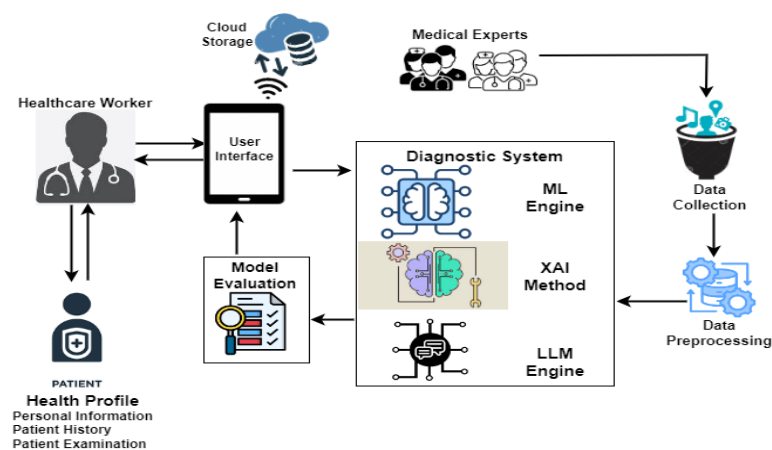


Figure 1. Malaria and Typhoid Fever Diagnosis Framework.

2.2. Description of the dataset used for the study

The New Frontiers in Research Fund project's dataset was used in this study and the dataset comprising 4870 patient records, was organized into six sections and included demographic data, patient symptoms, risk factors, and diagnosis information. The first section contains the patient demographics and the diagnosing physician's information. The second section, contains the patient's symptoms on a five-point scale (1=absent; 2=mild; 3=moderate; 4=severe; 5=very severe), along with the doctor's level of confidence (a numerical rating scale from 1 to 10). The patient's degree of susceptibility to the other non-clinical risk factors was listed in the third section, and the doctor's initial diagnosis was listed in the fourth. The confirmed diagnosis was included in the fifth section of the dataset and the further investigation conducted on the patient was reported in the sixth section. A linguistic scale (1=absent; 2=very low; 3=low; 4=moderate; 5=high; 6=very high) was used to rate the intensity of attack for both preliminary and confirmed diagnoses (Sections 4 and 5), along with the doctor's degree of confidence (1–10) in each case.

Table 1. Statistics of male and female patients in the dataset.

Patient Age (Years)	< 5	5 -12	13 – 19	20 – 64	≥ 65	Total
Male	534	346	150	1012	133	2175
Female	419	323	213	1605	135	2695
Total	953	669	363	2617	268	4870

2.3. Data Preprocessing and Oversampling

The collected dataset comprised columns with both numeric and string data types, along with a few missing values. Data preprocessing was done, including feature selection, feature scaling, and data cleaning. Records with missing features, irrelevant data, and columns that were not needed were eliminated during the data-cleaning process. Records with missing symptoms were likewise eliminated to maintain the integrity of the dataset. Because the patient consultation tool did not include symptoms for patients under the age of five, records of those patients were removed. This is because patients in this age group may not be able to accurately express certain symptoms, leaving them to rely entirely on their parents' interpretation. A selection of the most pertinent and significant features for modeling febrile illness (malaria and typhoid fever symptoms) was made to carry out the feature selection process. A patient's symptoms, the intensity of each symptom, and confirmed diagnoses (malaria and typhoid fever) are all included in the processed dataset. The dataset was reduced to 3914 records including only the malaria and typhoid fever confirmed diagnoses and their twelve (12) symptoms. The list of symptoms and diseases with abbreviations is presented in Table 2.

As shown in Figure 2, custom mapping was used to map the disease severity ‘Absent’ (1) to binary 0 and ‘Very-low’ to ‘Very-severe’ (2 to 6) to binary 1.

Table 2. Symptoms and diseases with abbreviations.

Symptom/Disease	Abbreviation
Abdominal pains	ABDPN
Bitter taste in mouth	BITAIM
Chills and rigors	CHLNRIG
Constipation	CNST
Fatigue	FTG
Fever	FVR
Generalized body pain	GENBDYPN
Headaches	HDACH
High-grade fever	HGGDFVR
Lethargy	LTG
Muscle and body pain	MSCBDYPN
Stepwise rise fever	SWRFVR
Malaria	MAL
Typhoid fever/ Enteric fever	ENFVR

	ABDPN	BITAIM	CHLNRIG	CNST	FTG	FVR	HGGDFVR	SWRFVR	GENBDYPN	HDACH	LTG	MSCBDYPN	MAL	ENFVR
0	4	4	4	3	4	5	4	3	4	3	4	4	1	1
1	2	2	2	1	2	2	3	4	1	1	1	1	1	0
2	1	3	5	3	3	5	4	3	4	4	4	4	0	1
3	3	3	1	1	1	3	1	3	4	3	1	2	1	0
4	3	1	1	1	3	4	4	3	4	4	1	3	0	1
...	...	...	...	...	...	...	...	...	...	...	...	...	...	...
3909	1	1	1	1	1	1	1	1	1	1	1	1	1	0
3910	1	1	1	1	4	1	1	1	1	4	1	1	1	0
3911	1	1	1	1	1	1	1	1	1	1	1	1	1	0
3912	1	1	1	1	1	1	1	1	1	1	1	1	0	0
3913	1	1	1	1	4	1	1	1	1	4	1	1	1	0

3914 rows × 14 columns

Figure 2. Pre-processed dataset.

After further analysis, we noticed that out of the 3914 patients, 1088 patients had neither malaria nor typhoid fever, 1669 had only malaria, 107 had only typhoid fever and 1050 had both diseases. Synthetic minority oversampling technique (SMOTE) was employed to handle the class imbalance. SMOTE has several advantages and when compared to just replicating minority class instances, it lowers the chance of overfitting by creating synthetic samples. It improves model performance, is compatible with most ML techniques, and is useful for various types of data. SMOTE identified minority class instances, selected k-nearest neighbors, generated and added synthetic samples to the original dataset as presented in Figure 3. The oversampled dataset contains 6676 patient records with the class labels 0(No disease), 1(Typhoid only), 2(Malaria only), and 3(Both diseases) in the ‘Disease’ column.

	ABDPN	BITAIM	CHLNRI	CNST	FTG	FVR	HGGDFVR	SWRFVR	GENBDYPN	HDACH	LTG	MSCBDYPN	disease
0	4	4	4	3	4	5	4	3	4	3	4	4	3
1	2	2	2	1	2	2	3	4	1	1	1	1	2
2	1	3	5	3	3	5	4	3	4	4	4	4	1
3	3	3	1	1	1	3	1	3	4	3	1	2	2
4	3	1	1	1	3	4	4	3	4	4	1	3	1
...	...	...	...	...	...	...	...	...	...	...	...	...	...
6671	2	3	3	1	2	2	3	1	3	4	3	3	3
6672	1	3	3	1	3	4	2	1	3	3	2	3	3
6673	4	1	1	1	3	1	1	1	2	2	3	2	3
6674	3	3	2	1	4	4	3	4	3	3	2	3	3
6675	4	1	1	1	1	2	1	1	2	1	1	1	3

6676 rows × 13 columns

Figure 3. Oversampled dataset with SMOTE.

2.4. Diagnostic Models and Model Optimization

We use Google Colaboratory (Colab), a free cloud-based platform from Google that offers a Python programming environment with quick access to robust graphics processing unit (GPU) resources and ML libraries. Additionally, the platform provides a CPU runtime and easily integrates Google Drive for storage. Python packages and libraries such as NumPy, Pandas, Scikit-Learn, and Matplotlib which are necessary for creating classification models were utilized. The ML algorithms used in building our diagnostic models and the performance metrics are presented in the subsection incorporating hyperparameter tuning called grid search cross-validation (GridSearchCV) to increase the precision of the diagnosis. GridSearchCV is an expanded method for optimizing hyperparameters by enabling customized search spaces for each hyperparameter, using designated ranges. Our local laptops utilized for this study were Dell Latitude 7480, Core i5-7200U CPU @ 2.50GHz (4 CPUs), ~2.7GHz with 16GB RAM for the ML and XAI modeling while Samsun 950QDB, Core i7-1165G7 @ 2.80Ghz (8 CPUs) ~ 2.8GHz with 16GB RAM was used for the large language modeling. We used Visual Studio Code, a free coding editor that supports several extensions and allows for quick coding initiation. Large language models are easily accessible thanks to the Python foundation of our development environment. The process was automated by utilizing core Python packages and libraries, such as NumPy, Pandas, os, json, re, pyperclip, and io. The prompt generator, automator, and interpreter are the components of the information extractor. The prompt generator converts data into a readable format, saves prompts into a JSON file, and divides the patient's symptoms and severity into manageable chunks for prompts. The prompt used by Caruccio et al. [23] mimics the conversation a physician when seeking assistance in diagnosing a patient based on particular symptoms. The template is “The patient has these symptoms: [S] Tell me which of the following diagnoses is most related to the symptoms: [D]? [H]”. This template was modified to arrive at our prompt: “The patient has these symptoms with severity levels, listed in the table below. (create a table with only the diagnosis column filled in), the output should be in CSV format, diagnosis [Malaria, Typhoid Fever, Both, None]?”. The automator manages data flow by retrieving outputs and storing them in a JSON file. It feeds these prompts into the large language models (GPT, Gemini, and Perplexity). After that, the interpreter transforms the JSON output into an Excel file so that reporting and analysis of the data can be carried out. The link (https://github.com/FebrileDiseasesDiagnoses/Auto_tool.git) to the scripts can be found in this GitHub account.

2.4.1. Random Forest

Random Forest algorithm is widely employed for tasks involving both regression and classification. It is an ensemble learning technique that combines several decision trees to increase prediction accuracy as shown in Figure 4. Due to its simplicity, RF is arguably the most used algorithm and usually produces excellent results even without hyperparameter tuning. According to Han et al. [24], RF is one of the most sophisticated ensemble algorithms with a robust resistance to over-fitting. It is more accurate, trains predictions concurrently, operates well on big datasets, and is very good at estimating missing data [25]. RF is helpful in pattern recognition for resolving high-dimensional and complex problems, including the prediction of disease conditions [26–28]. By combining individual tree predictions via voting (classification), the final prediction is produced. This ensemble approach increases the model's robustness, decreases overfitting, and can aid in diagnosing febrile diseases.

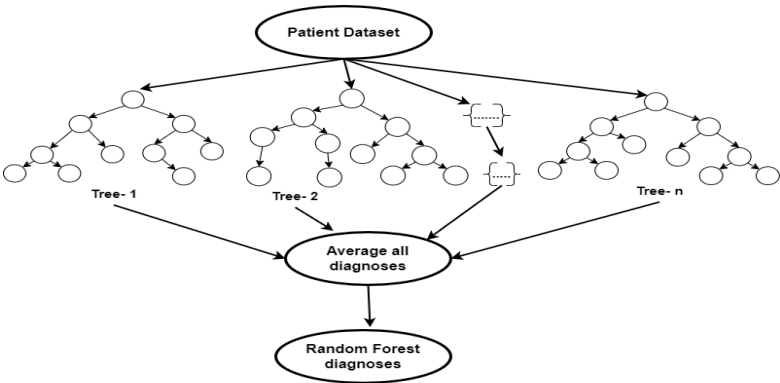


Figure 4. Random Forest schematic diagram.

2.4.2. Extreme Gradient Boost

XGBoost is a machine learning algorithm that is a component of the gradient boosting framework, which is a subset of ensemble learning that can be applied to regression or classification predictive modeling issues. Figure 5 depicts the computation process used by XGBoost. Its foundation is gradient-boosted decision trees and regularization techniques for enhancing model generalization. XGBoost introduces weak learners into the ensemble, focusing each new learner on correcting the mistakes made by the previous ones. Because of its reputation for managing structured data, XGBoost is extensively utilized in numerous applications, including the prediction of disease [32].

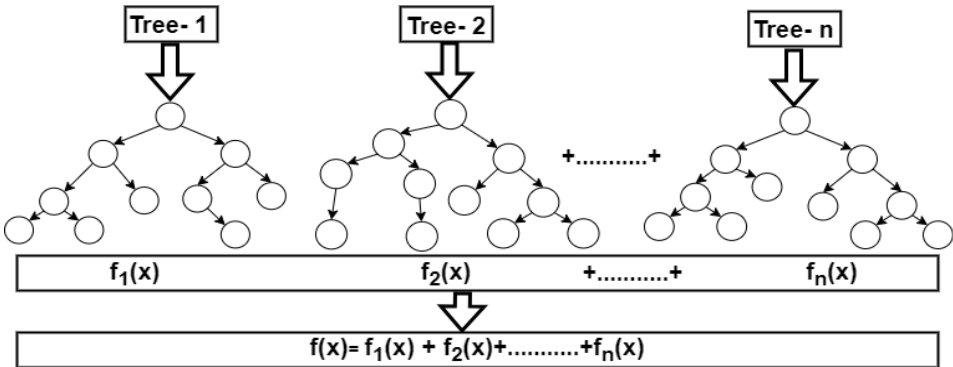


Figure 5. Extreme gradient boosting schematic diagram.

2.4.3. Support Vector Machine

SVM is broadly used for regression and classification tasks. SVM seeks to determine which hyperplane best divides the data into distinct classes. The margin is the distance between the hyperplane and the closest observations, and the support vectors are the points that are closest to it as shown in Figure 6. In particular, SVM is well-known for working well in high-dimensional spaces

and for handling non-linearly separable data by utilizing kernel functions. SVM uses little memory, performs well with a wide range of features, and can be tailored with various kernel functions for intricate decision boundaries. SVM can handle high-dimensional data and is resistant to overfitting as well as binary and multi-class classification issues in medical diagnosis, making it an effective tool for diagnosing diseases [33].

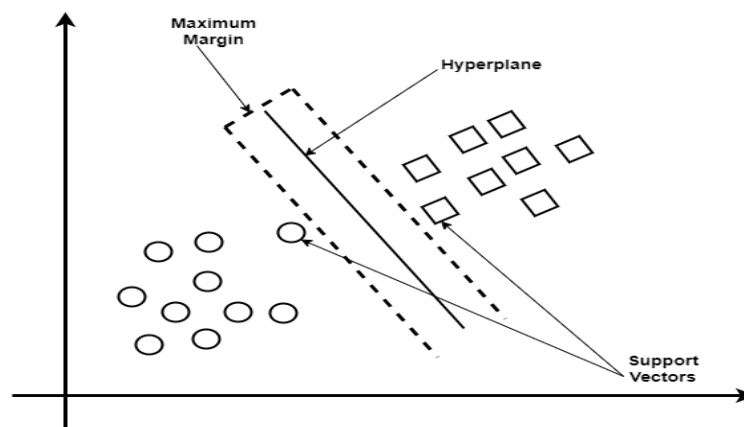


Figure 6. Support Vector Machine diagram.

2.5. Interpretability and Explainability Methods

Local Interpretable Model-agnostic Explanations (LIME)- approximates the complex model near a particular prediction with an interpretable model such as a linear model to provide local explanations. The following steps are applied for LIME Integration: (i) Instance Selection- for each model, LIME was applied to every instance of the test dataset to provide detailed explanations of diagnoses. (ii) Feature Contributions- LIME generates visual explanations that show how these features influenced the diagnosis and features with bars to the right of the plot increase the probability of a certain disease, while those to the left decrease it. (iii) Aggregation for Global Insight- By averaging LIME explanations across several cases, significant insights into which features for each disease diagnosis can be learned.

Generative Pre-trained Transformer (GPT)- is pre-trained using unsupervised learning on a large corpus of text data, where it learns to predict the word that will appear next in a sequence based on every word that has come before it. This pre-training gives GPT a thorough grasp of language syntax, semantics, and context. When GPT is fine-tuning on particular tasks, like text generation, question answering, or text completion, it uses its learned representations to produce outputs that make sense and are relevant to the context. GPT is an effective tool for natural language processing applications because it can produce text similar to that of a human being and manage a wide range of linguistic tasks.

Bidirectional Encoder Representations from Transformers (BERT)- is a transformer-based model that is trained to predict missing words in both directions with the help of masking certain words in the input and making predictions about them using both left and right context. Thanks to this bidirectional training, it can capture more complex contextual meanings and relationships within text, producing more accurate language representations. BERT can comprehend subtleties in language and perform well on a variety of natural language understanding tasks, including named entity recognition, sentiment analysis, and machine translation, thanks to its large-scale pre-training tasks. BERT is a flexible and potent model for a range of natural language processing applications. Its performance can be further improved by fine-tuning it for particular tasks.

2.6. Model Performance metrics

The experimental models were evaluated using key performance metric components. When evaluating the sensitivity and specificity of diagnostic tests, these metric components are helpful. The evaluation metrics listed below were used in this paper.

Accuracy- a measurement of how well a model predicts all labels linked to each data point in a dataset. Datasets with a balanced distribution of positive and negative samples are good candidates for accuracy. For unbalanced datasets, it is less helpful because it can be deceptive.

$$Accuracy = \frac{True\ Positives + True\ Negatives}{True\ Positives + True\ Negatives + False\ Positives + False\ Negatives} \quad (1)$$

Precision- a metric that expresses how accurately a model predicts positive outcomes; it measures the model's capacity to correctly identify true positive instances while avoiding false positives. When false positives come at a high cost, accuracy matters. In the context of medical diagnosis, for instance, a false positive may result in needless treatments.

$$Precision = \frac{True\ Positives}{True\ Positives + False\ Positives} \quad (2)$$

Recall- a metric used to assess a model's capacity to locate every positive instance in a dataset. It measures how sensitive the model is to True Positive cases. When the cost of false negatives is significant, as in medical screenings, recall is crucial because it can be crucial to miss a positive case (false negative).

$$Recall = \frac{True\ Positives}{True\ Positives + False\ Negatives} \quad (3)$$

F1-Score- a metric that represents the harmonic mean of recall and precision. The F1-score is limited to a range of binary values, where 1 denotes that every class's data point was correctly predicted and 0 denotes that any class's data point was incorrectly predicted. When you must strike a balance between recall and precision, the F1 Score can be helpful, particularly when your class distribution is not uniform.

$$F1 = 2 * \left(\frac{Precision * Recall}{Precision + Recall} \right) \quad (4)$$

Log Loss- a measure of the probability of a prediction's accuracy and it penalizes the difference between the expected probabilities and the actual class labels. Log loss is helpful when one needs a metric that can handle probabilistic model outputs and penalizes incorrect predictions more severely.

$$Log\ loss = -\frac{1}{N} \sum_{i=1}^N [y_i \log(p_i) + (1 - y_i) \log(1 - p_i)] \quad (5)$$

Where N is the total number of samples in the dataset. y_i is the actual label of the i -th instance. p_i is the predicted probability of the i -th instance being in the positive class. $\log(p_i)$ is the natural logarithm of the predicted probability for the positive class

3. Results and Discussion

The results of our assessment of the models' performance are shown in this section, the XAI method adopted as well as the experimental assessment of the large language models of febrile disease diagnoses (Malaria and Typhoid Fever). Figures 7–9 present the confusion matrices, an essential instrument for assessing how well a classification ML model performs.

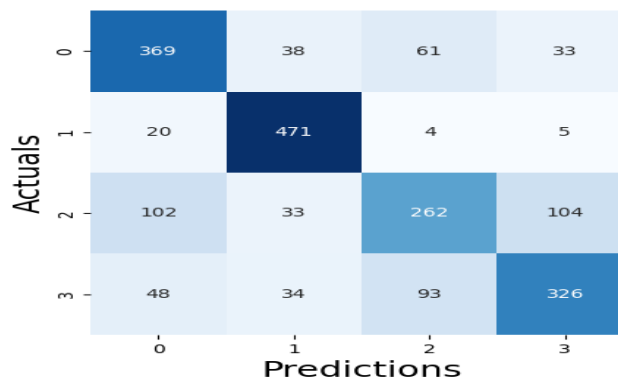


Figure 7. XGBoost Algorithm Confusion Matrix.

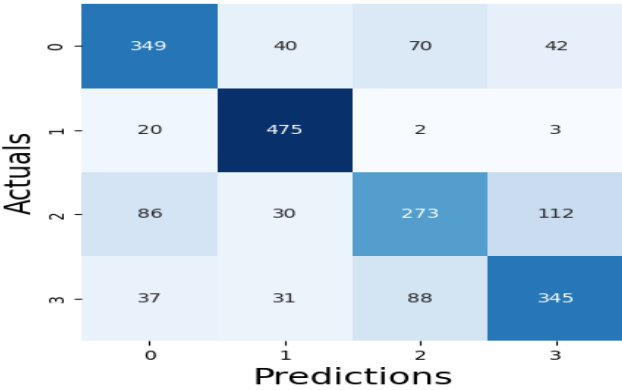


Figure 8. RF Algorithm Confusion Matrix.

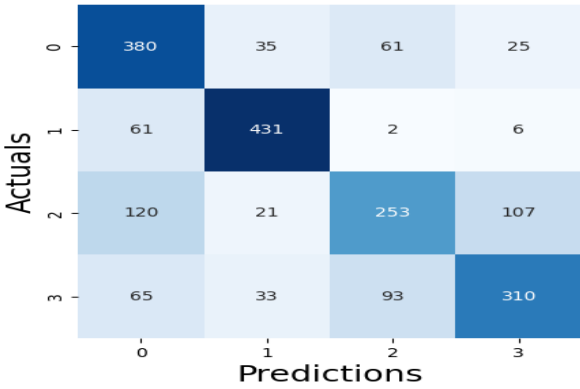


Figure 9. SVM Algorithm Confusion Matrix.

Table 3 presents values of these metrics and the computation time of each model while Figure 10 is a pictorial representation of the model’s performance based on the considered metrics. The result shows that RF (accuracy = 71.99%, precision =71.29%, recall=71.99%, F1-Score=71.45%) demonstrates superior performance, outperforming XGBoost (accuracy=71.29%, precision=70.56%, recall=71.29%, F1-Score=70.66%) and SVM (accuracy= 68.60%, precision=68.65%, recall=68.60%, F1-Score=68.21%). High recall and precision are essential for diagnosing diseases like typhoid and malaria by guaranteeing that the majority of real cases are identified. In this case, high precision helps prevent needless treatments for illnesses that are not present. Because both XGBoost and RF do a good job of balancing these metrics, they are better suited for clinical applications where false positives and false negatives can have detrimental effects. Also, XGBoost has a smaller log loss, which suggests more accurate and well-calibrated probability estimates as well as stronger diagnosis confidence and this may be critical in medical diagnostics, where accuracy is not as important as confidence in the presence of a disease. In medical scenarios where treatment decisions are influenced by the certainty of a diagnosis, lower log loss values for XGBoost indicate that its probability estimates are more reliable. Because of RF's higher log loss, probability estimates are less trustworthy, which could cause uncertainty when making decisions. SVM performs worse than the other two models in terms of performance metrics and computation time (running time exceeds one hour), implying that it might not work as well for diagnosing typhoid and malaria in this specific dataset. Therefore, ensemble techniques (XGBoost and Random Forest) may be better at capturing the intricate relationships between symptoms and diseases than the SVM model.

Table 3. Diagnostics model performance.

Algorithm	Accuracy	Precision	Recall	F1-score	Log Loss	Computation time
XGBoost	0.7129	0.7056	0.7129	0.7066	0.7808	2 minutes, 32 seconds

RF	0.7199	0.7129	0.7199	0.7145	1.0548	14 minutes, 9 seconds
SVM	0.6860	0.6865	0.6860	0.6821	0.8016	1hr, 22 minutes, 7 seconds

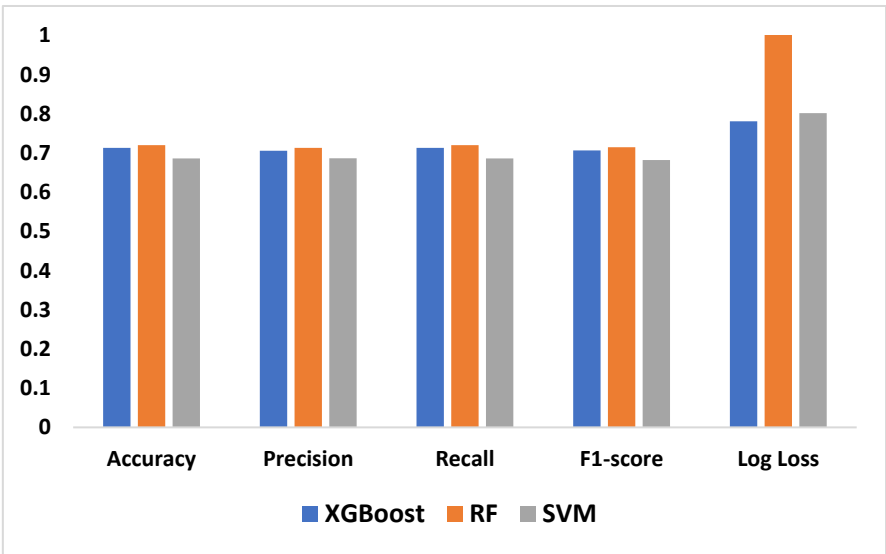


Figure 10. Performance Evaluation of the Machine Learning Models.

The LIME plots (Figures 11–13) provide a global view of how the features(symptoms) contribute to the model’s diagnoses across the entire test dataset, identifying features with the highest average contributions, both positively and negatively across all diagnoses. The XGBoost LIME diagram in Figure 11 shows that symptoms such as SWRFVR, HDACH, and CNST as specified by their negative contributions on the left side of the plot suggest that the lower levels or absence of these symptoms are associated with a lower likelihood of a patient having malaria and typhoid. Whereas symptoms such as BITAIM, LTG, CHLNRIG, MSCBDYPN, and FVR are the most influential symptoms constantly contributing to the diagnoses of malaria and typhoid across numerous patients.

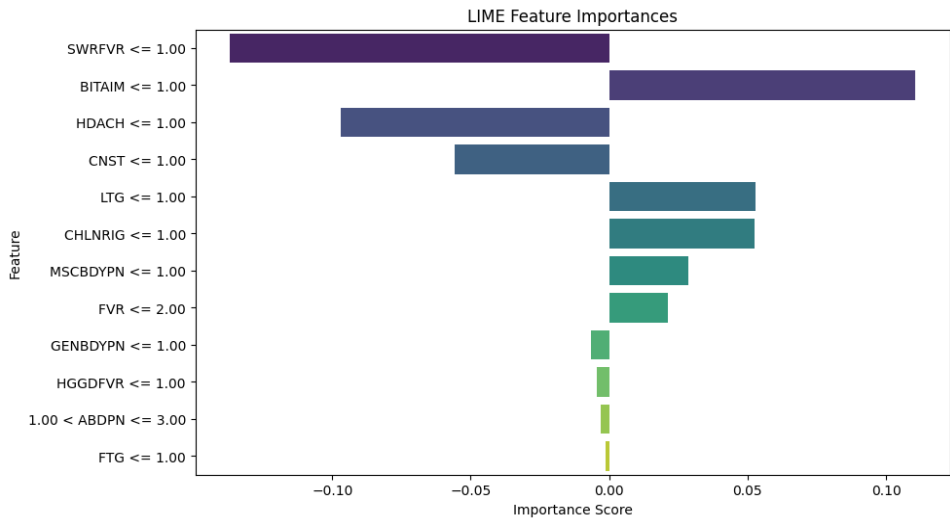


Figure 11. XGBoost Algorithm LIME diagram.

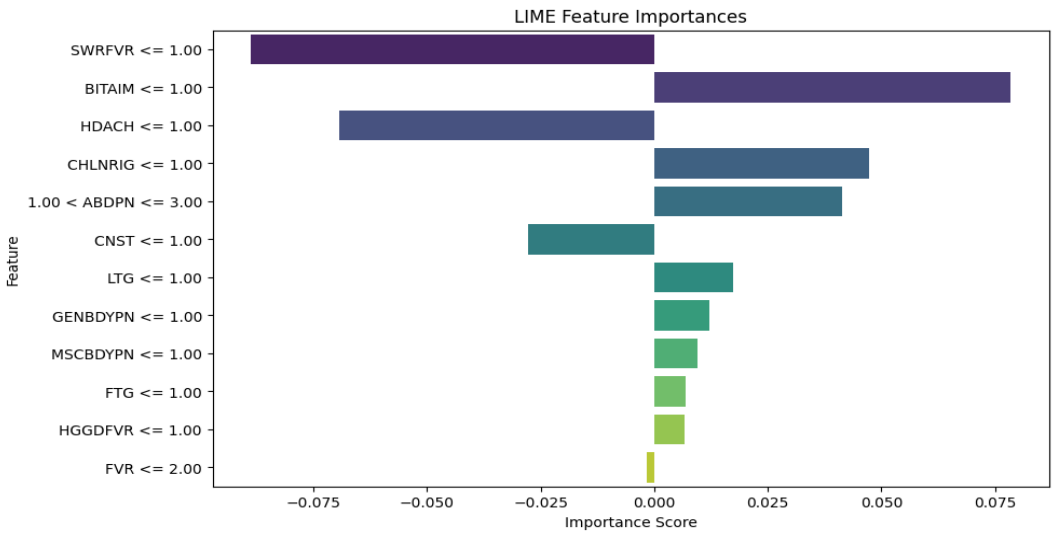


Figure 12. RF Algorithm LIME diagram.

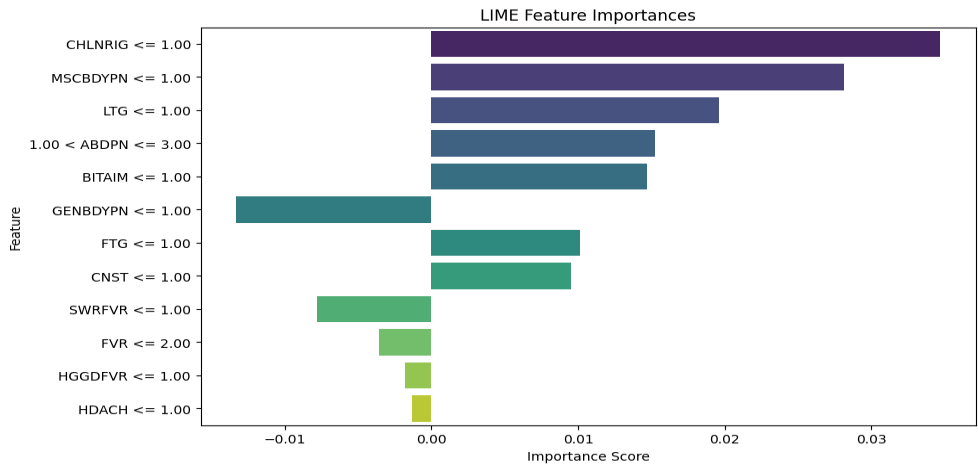


Figure 13. SVM Algorithm LIME diagram.

The RF LIME diagram in Figure 12 also points out that the same symptoms (SWRFVR, HDACH, and CNST) are associated with a lower likelihood of having malaria and typhoid whereas BITAIM, CHLNRIG, ABDPN, LTG, GENBDYPN, MSCBDYPN, FTG and HGGDFVR are influential symptoms that contribute to the diagnoses of malaria and typhoid among patients.

Figure 13 shows the SVM LIME diagram, indicating CHLNRIG has the highest feature importance, followed by MSCBDYPN, LTG, ABDPN, BITAIM, FTG, and CNST as the influential symptoms that contribute to the diagnoses of malaria and typhoid among patients while GENBDYPN, SWRFVR, FVR, HGGDFVR and HDACH are associated with a lower likelihood of having malaria and typhoid.

It is observed that medical experts should focus on these influential symptoms for the diagnosis of malaria and typhoid fever in patients. They are BITAIM, CHLNRIG, LTG, ABDPN, MSCBDYPN, FVR, GENBDYPN, FTG, and HGGDFVR. This is consistent with the results of Asuquo et al. [34] where GENBDYPN, CHLNRIG, ABPN, FVR, FTG, and HGDFVR were observable symptoms. LIME has numerous advantages. It explains the individual diagnosis in a form that is relatively easy for humans to comprehend, aiding healthcare workers to understand why a model made a specific diagnosis. LIME can be applied to many ML models and this versatility makes it suitable for various medical diagnostic systems. Besides, LIME is suitable for generating explanations using local approximations [35]. The limitation of LIME is that it is computationally intensive and expensive to generate explanations for individual diagnoses, especially for large datasets and complex models.

Furthermore, three sets of experiments were conducted to evaluate the performance of ChatGPT, Gemini, and Perplexity in diagnosing malaria and typhoid. In Experiment 1, one prompt was sent at a time to the LLMs for the first 100 patients in the dataset, recording the outputs in a CSV format to see how they performed with a single set of prompts. Experiment 2 involved sending 100 prompts from the first 100 patients in the dataset to the LLMs, and storing the outputs in a CSV format to observe their responses to a series of prompts. In Experiment 3, 100 unique prompts were sent to the models repeatedly until the entire dataset was exhausted, to assess how the models performed when given large sets of unique prompts. Table 4 presents the results of the three experiments. In Exp 1, ChatGPT 3.5 has a slightly better performance with the highest F1-score (30.99%) because the F1-score is crucial as it balances recall and precision, providing a comprehensive measure of the model's performance. Although better accuracy and recall are achieved by ChatGPT 3.5 and Gemini (30%), Perplexity is better at minimizing false positives with its highest precision (38.90%). In Exp 2, Perplexity performs better, with the highest F1-score (26.29%), accuracy (28%), and recall (28%). Because it provides a comprehensive measure of the model's performance by balancing recall and precision, the F1-score is especially significant. ChatGPT 3.5 is better at reducing false positives with the highest precision, while Gemini has the lowest performance. In Exp 3, ChatGPT 3.5 has better accuracy, precision, and recall followed by Gemini and Perplexity. Although the ChatGPT model may be having trouble striking a balance between minimizing false positives and identifying true positives, the model's relatively low F1- score suggests that there may be an imbalance between precision and recall.

Table 4. Large language models' performance.

Experiment	Algorithm	Accuracy	Precision	Recall	F1-score
Exp 1	Chat GPT 3.5	0.3000	0.35562	0.3000	0.30999
	Gemini	0.3000	0.3449	0.3000	0.2908
	Perplexity	0.2600	0.3890	0.2600	0.28736
Exp 2	Chat GPT 3.5	0.2600	0.2909	0.2600	0.2615
	Gemini	0.2700	0.2607	0.2700	0.2296
	Perplexity	0.2800	0.2524	0.2800	0.2629
Exp 3	Chat GPT 3.5	0.3297	0.3324	0.3297	0.2926
	Gemini	0.2895	0.2709	0.2895	0.2728
	Perplexity	0.2632	0.1957	0.2632	0.2171

ChatGPT is an innovative technological tool for comprehending and processing natural language, making it suitable for interpreting and summarizing complex up-to-date information. Gemini is an adaptable tool that can handle various data types such as images and text, making it suitable for diagnostic purposes. Perplexity is specialized in comprehending and generating complex queries as well as maintaining context that can be vital for the retrieval and analysis of medical research. These LLMs lack specialized knowledge and are also capable of producing inaccurate answers which can be critical in a medical context. They require high computational power to generate and process responses which could limit real-time systems. Data security and patient privacy are concerns when handling sensitive medical data and they require proper validation and regulatory approval before they can be trusted and adopted for clinical use. To facilitate healthcare professionals' comprehension of the reasoning behind a diagnostic output, LLMs integrate and analyze large amounts of medical data and produce human-readable explanations for their decisions.

The overall ML models' performance in the study was moderate, suggesting the need for a sufficient dataset to enhance the diagnostic models. While the traditional SMOTE aided in balancing

the dataset, employing an advanced oversampling method may help in improving the model performance. Even with GridSearchCV, the hyperparameters might still be improved, particularly for SVM. Better configurations could result from investigating alternative parameter tuning techniques like RandomizedSearchCV or Bayesian Optimization. To improve the results of the LLMs, the LLMs will be fine-tuned with a larger dataset, and an ensemble method will be employed to combine the strengths of different LLMs.

To integrate ML, XAI, and LLM techniques into an app, we propose two methods.

Method 1: Separate Training and Validation for ML and LLM

1. Train, test, and validate an ML model to diagnose malaria and typhoid based on the patient dataset
2. Apply LIME to explain the ML models' diagnoses and how each symptom contributed to the diagnoses
3. Train, Test, and Validate an LLM model independently for generating explanations based on the patient dataset
4. Integrate the outputs from ML, LIME, and LLM to provide a comprehensive and interpretable diagnosis.

The advantage of method 1 is that it might lead to higher diagnostic performance considering the specific training of the two models (ML and LLM) for this task. The disadvantage is that the training and validation process of two independent models would increase the computational complexity of the diagnostic system, especially in combining the results to ensure consistency and coherence.

Method 2: Integrated ML, LIME, and LLM Process

1. Train, Test, and Validate an ML model to diagnose malaria and typhoid based on the patient dataset
2. Apply LIME to explain the ML models' diagnoses and how each symptom contributed to the diagnoses
3. Use LLM for further explainability by passing the patient symptoms and ML results (with LIME explanations) through the model to generate diagnostic explanations in natural language.

The advantage of method 2 is its simplicity because an integrated pipeline reduces complexity, making the system easier to develop, test, and maintain which we recommend for implementation. Performance will be increased and computational overhead could be decreased by streamlining the procedure into a single pipeline. The explanations produced by LIME can be directly considered by the LLM, which may result in more logical and contextually appropriate explanations. The disadvantage is that the quality of the initial ML and LIME outputs determines the quality of the explanations provided by the LLM.

Whereas previous studies [36–39] applied ML models for disease diagnoses, this paper integrated the use of XAI and LLMs, to enhance transparency and interpretability in the diagnostic processes. The use of LIME for feature importance analyses and LLMs for generating context-aware explanations have distinguished the present study. Several factors can contribute to the low performance scores in Table 4. These include: 1) the dataset used during the training is limited in size and diversity, affecting the models' ability to generalize to unseen cases. 2) LLMs may require further fine-tuning and optimization, as the complexity of the diseases being diagnosed may overlap with other illnesses thereby challenging the models to accurately differentiate between them. Furthermore, LLMs did not show high domain tolerance to the investigated illnesses, hence, fine-tuning them on domain-specific data can significantly improve their performance.

4. Conclusions

This study creates a medical diagnostic framework for Malaria and Typhoid fever by integrating XAI, LLMs, and ML models. This approach aims to demystify the black-box nature of ML models, offering transparent insights into how each feature or symptom affects the diagnosis. The RF model showed superior prediction performance across all metrics compared to XGBoost and SVM. The high

recall and precision values in RF are crucial for accurately diagnosing these diseases, and preventing unnecessary treatments. However, XGBoost exhibited the lowest log loss (0.7808) and fastest computation time, indicating more reliable probability estimates and stronger diagnosis confidence, which is vital for treatment decisions. Further analysis indicates that SVM performs worse than the other two models in terms of performance metrics and computation time, making it less suitable for this dataset. The study suggests that ensemble techniques like RF and XGBoost better capture the complex relationships between symptoms and diseases. The XAI analysis identified BITAIM, CHNLNRIG, LTG, ABDPN, MSCBDYPN, FVR, GENBDYPN, FTG, and HGGDFVR as key features for predicting Malaria and Typhoid. Among LLMs, ChatGPT 3.5 performed slightly better than Gemini and Perplexity. This study recommends integrating ML, LIME, and LLM processes due to the simplified workflow of the overall development and maintenance process, resource efficiency, and improved explainability as a result of passing both patient symptoms and ML results through the LLM as the LLM can take into account the full context provided by LIME.

Author Contributions: Conceptualization, F.-M.U., and K.A.; methodology, K. A., C.A., D. A., and M.E.; validation, F.-M.U., O.O., K. A., C.A., D.A, and M.E.; formal analysis, K. A., P.A., and D. A.; data curation, K.A., and P. A.; writing—original draft preparation, K. A., D.A., P. A., E. J. A. J. and M.E.; writing—review and editing, K. A., M. E., D. A., O.O., C. A., O.M. and F.-M.U. supervision, F.-M.U. C. A., D. A., O.O., and M.E.; project administration, F.-M.U., and O.O.; funding acquisition, F.-M.U. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by the New Frontier Research Fund, grant number NFRFE-2019-01365 between April and March 2024.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: Data will be made available on reasonable request.

Conflicts of Interest: The authors declare no conflict of interest.

References

- Asuquo, D.; Attai, K.; Obot, O.; Ekpenyong, M.; Akwaowo, C.; Arnold, K.; Uzoka, F. M. Febrile disease modeling and diagnosis system for optimizing medical decisions in resource-scarce settings. *Clin. eHealth* **2024**, *7*, 52-76. DOI: 10.1016/j.ceh.2024.05.001
- Galán, J. E. Typhoid toxin provides a window into typhoid fever and the biology of *Salmonella* Typhi. *Proc. Natl. Acad. Sci. U. S. A.* **2016**, *113* (23), pp. 6338-6344. DOI: 10.1073/pnas.1606335113
- Gashaw, T.; Jambo, A. Typhoid in less developed countries: a major public health concern. In *Hygiene and Health in Developing Countries-Recent Advances*; IntechOpen **2022**. DOI: 10.5772/intechopen.108109
- Alhumaid, N. K.; Alajmi, A. M.; Alosaimi, N. F.; Alotaibi, M.; Almangour, T. A.; Nassar, M. S.; Tawfik, E. A. Reported Bacterial Infectious Diseases in Saudi Arabia. Overview and Recent Advances **2023**, pp. 1-39. DOI:10.21203/rs.3.rs-3351846/v1
- Paton, D. G.; Childs, L. M.; Itoe, M. A.; Holmdahl, I. E.; Buckee, C. O.; Catteruccia, F. Exposing *Anopheles* mosquitoes to antimalarials blocks *Plasmodium* parasite transmission. *Nature* **2019**, *567* (7747), pp. 239-243. DOI:10.1038/s41586-019-0973-1
- Sato, S. *Plasmodium*—a brief introduction to the parasites causing human malaria and their basic biology. *J. Physiol. Anthropol.* **2021**, *40* (1), 1. DOI: 10.1186/s40101-020-00251-9
- Carson, B. B., III. Mosquitos and Malaria Take a Toll. In *Challenging Malaria: The Private and Social Incentives of Mosquito Control*; Springer International Publishing: Cham **2023**, pp. 15-25. DOI:10.1007/978-3-031-39510-9
- Attai, K.; Amannejad, Y.; Vahdat Pour, M.; Obot, O.; Uzoka, F. M. A systematic review of applications of machine learning and other soft computing techniques for the diagnosis of tropical diseases. *Trop. Med. Infect. Dis.* **2022**, *7* (12), 398. DOI: 10.3390/tropicalmed7120398
- Boina, R.; Ganage, D.; Chincholkar, Y. D.; Wagh, S.; Shah, D. U.; Chinthamu, N.; Shrivastava, A. Enhancing Intelligence Diagnostic Accuracy Based on Machine Learning Disease Classification. *Int. J. Intell. Syst. Appl. Eng.* **2023**, *11* (6s), pp. 765-774.
- Asuquo, D. E.; Umoren, I.; Osang, F.; Attai, K. A Machine Learning Framework for Length of Stay Minimization in Healthcare Emergency Department. *Stud. Eng. Technol. J.* **2023**, *10* (1), 1-17. DOI:10.11114/set.v10i1.6372.

11. Anderson, J.; Thomas, J. Interpretable Machine Learning Models for Healthcare Applications. *EasyChair* **2024**, No. 12358.
12. Albahri, A. S.; Duhaim, A. M.; Fadhel, M. A.; Alnoor, A.; Baqer, N. S.; Alzubaidi, L.; Deveci, M. A systematic review of trustworthy and explainable artificial intelligence in healthcare: Assessment of quality, bias risk, and data fusion. *Inf. Fusion* **2023**. DOI:10.1016/j.inffus.2023.03.008
13. Tan, L.; Huang, C.; Yao, X. A Concept-Based Local Interpretable Model-Agnostic Explanation Approach for Deep Neural Networks in Image Classification. In *International Conference on Intelligent Information Processing*; Springer Nature Switzerland: Cham **2024**; pp. 119-133. DOI: 10.1007/978-3-031-57919-6_9
14. Thombre, A. Comparison of decision trees with Local Interpretable Model-Agnostic Explanations (LIME) technique and multi-linear regression for explaining support vector regression model in terms of root mean square error (RMSE) values. *arXiv Preprint* **2024**, arXiv:2404.07046. DOI:10.48550/arXiv.2404.07046
15. Okay, F. Y.; Yıldırım, M.; Özdemir, S. Interpretable machine learning: A case study of healthcare. In *2021 International Symposium on Networks, Computers and Communications (ISNCC)*; IEEE: October **2021**; pp. 1-6. DOI: 10.1109/ISNCC52172.2021.9615727
16. Attai, K.; Akwaowo, C.; Asuquo, D.; Esubok, N. E.; Nelson, U. A.; Dan, E.; Uzoka, F. M. Explainable AI modelling of Comorbidity in Pregnant Women and Children with Tropical Febrile Conditions. In *International Conference on Artificial Intelligence and its Applications*; December **2023**, pp. 152-159. DOI: 10.59200/ICARTI.2023.022
17. Ashraf, K.; Nawar, S.; Hosen, M. H.; Islam, M. T.; Uddin, M. N. Beyond the Black Box: Employing LIME and SHAP for Transparent Health Predictions with Machine Learning Models. In *2024 International Conference on Advances in Computing, Communication, Electrical, and Smart Systems (iCACCESS)*; IEEE: March **2024**; pp. 1-6. DOI: 10.1109/iCACCESS61735.2024.10499522
18. Li, F.; Jin, Y.; Liu, W.; Rawat, B. P. S.; Cai, P.; Yu, H. Fine-tuning bidirectional encoder representations from transformers (BERT)-based models on large-scale electronic health record notes: an empirical study. *JMIR Med. Inform.* **2019**, *7* (3), e14830. DOI:10.2196/14830
19. Nakamura, Y.; Hanaoka, S.; Nomura, Y.; Nakao, T.; Miki, S.; Watadani, T.; Abe, O. Automatic detection of actionable radiology reports using bidirectional encoder representations from transformers. *BMC Med. Inform. Decis. Mak.* **2021**, *21*, pp. 1-19. DOI: 10.1186/s12911-021-01623-6
20. Gorenstein, L.; Konen, E.; Green, M.; Klang, E. BERT in radiology: a systematic review of natural language processing applications. *J. Am. Coll. Radiol.* **2024**. DOI: 10.1016/j.jacr.2024.01.012
21. Yenduri, G.; Ramalingam, M.; Selvi, G. C.; Supriya, Y.; Srivastava, G.; Maddikunta, P. K. R.; Gadekallu, T. R. GPT (generative pre-trained transformer)-a comprehensive review on enabling technologies, potential applications, emerging challenges, and future directions. *IEEE Access* **2024**. DOI: 10.48550/arXiv.2305.10435
22. Wang, Z.; Guo, R.; Sun, P.; Qian, L.; Hu, X. Enhancing Diagnostic Accuracy and Efficiency with GPT-4-Generated Structured Reports: A Comprehensive Study. *J. Med. Biol. Eng.* **2024**, 1-10. DOI: 10.1007/s40846-024-00849-9
23. Caruccio, L.; Cirillo, S.; Polese, G.; Solimando, G.; Sundaramurthy, S.; Tortora, G. Can ChatGPT provide intelligent diagnoses? A comparative study between predictive models and ChatGPT to define a new medical diagnostic bot. *Expert Syst. Appl.* **2024**, *235*, 121186. DOI: 10.1016/j.eswa.2023.121186
24. Han, H.; Zhang, Z.; Cui, X.; Meng, Q. Ensemble learning with member optimization for fault diagnosis of a building energy system. *Energy Build.* **2020**, *226*, 110351. DOI: 10.1016/j.enbuild.2020.110351
25. Zhu, L.; Qiu, D.; Ergu, D.; Ying, C.; Liu, K. A study on predicting loan default based on the random forest algorithm. *Procedia Comput. Sci.* **2019**, *162*, pp. 503-513. <https://doi.org/10.1016/j.procs.2019.12.017>
26. Ghosh, D.; Cabrera, J. Enriched random forest for high dimensional genomic data. *IEEE/ACM transactions on computational biology and bioinformatics* **2021**, *19*(5), pp. 2817-2828. DOI: 10.1109/TCBB.2021.3089417
27. Jackins, V.; Vimal, S.; Kaliappan, M.; Lee, M. Y. AI-based smart prediction of clinical disease using random forest classifier and Naive Bayes. *J. Supercomput.* **2021**, *77* (5), pp. 5198-5219. DOI:10.1007/s11227-020-03481-x
28. Palimkar, P.; Shaw, R. N.; Ghosh, A. Machine learning technique to prognosis diabetes disease: Random forest classifier approach. In *Advanced Computing and Intelligent Technologies: Proceedings of ICACIT 2021*; Springer Singapore, **2022**; pp. 219-244. DOI: 10.1007/978-981-16-2164-2_19
29. Kavzoglu, T.; Teke, A. Advanced hyperparameter optimization for improved spatial prediction of shallow landslides using extreme gradient boosting (XGBoost). *Bull. Eng. Geol. Environ.* **2022**, *81* (5), 201. DOI:10.1007/s10064-022-02708-w
30. Budholiya, K.; Shrivastava, S. K.; Sharma, V. An optimized XGBoost based diagnostic system for effective prediction of heart disease. *J. King Saud Univ.-Comput. Inf. Sci.* **2022**, *34* (7), pp. 4514-4523. DOI: 10.1016/j.jksuci.2020.10.013
31. Zhao, Y.; Li, X.; Li, S.; Dong, M.; Yu, H.; Zhang, M.; Gao, Z. Using machine learning techniques to develop risk prediction models for the risk of incident diabetic retinopathy among patients with type 2 diabetes mellitus: a cohort study. *Front. Endocrinol.* **2022**, *13*, 876559. DOI:10.3389/fendo.2022.876559

32. Asselman, A.; Khaldi, M.; Aammou, S. Enhancing the prediction of student performance based on the machine learning XGBoost algorithm. *Interact. Learn. Environ.* **2023**, *31* (6), 3360-3379. DOI:10.1080/10494820.2021.1928235
33. Devikanniga, D.; Ramu, A.; Haldorai, A. Efficient diagnosis of liver disease using support vector machine optimized with crows search algorithm. *EAI Endorsed Trans. Energy Web* **2020**, *7* (29);10. DOI:10.4108/eai.13-7-2018.164177
34. Asuquo, D. E.; Attai, K. F.; Johnson, E. A.; Obot, O. U.; Adeoye, O. S.; Akwaowo, C. D.; Uzoka, F. M. E. Multi-criteria decision analysis method for differential diagnosis of tropical febrile diseases. *Health Inform. J.* **2024**, *30* (2). DOI:10.1177/14604582241260659
35. Salih, A. M.; Raisi-Estabragh, Z.; Galazzo, I. B.; Radeva, P.; Petersen, S. E.; Lekadir, K.; Menegaz, G. A Perspective on Explainable Artificial Intelligence Methods: SHAP and LIME. *Adv. Intell. Syst.* **2024**, 2400304. DOI: 10.48550/arXiv.2305.02012
36. Maidabara, A. H.; Ahmadu, A. S.; Malgwi, Y. M.; Ibrahim, D. Expert system for diagnosis of malaria and typhoid. *Comput. Sci. IT Res. J.* **2021**, *2* (1); pp. 1-15. DOI: 10.51594/csitj.v2i1.274
37. Awotunde, J. B.; Imoize, A. L.; Salako, D. P.; Farhaoui, Y. An Enhanced Medical Diagnosis System for Malaria and Typhoid Fever Using Genetic Neuro-Fuzzy System. In *The International Conference on Artificial Intelligence and Smart Environment*; Springer International Publishing: Cham, **2022**; pp. 173-183. DOI: 10.1007/978-3-031-26254-8_25
38. Mariki, M.; Mkoba, E.; Mduma, N. Combining clinical symptoms and patient features for malaria diagnosis: machine learning approach. *Appl. Artif. Intell.* **2022**, *36* (1), 2031826. DOI: 10.1080/08839514.2022.2031826

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.