

A Transparent and Computationally Efficient AI Framework for Breast Cancer Diagnosis Using Predictive Modeling and Explainability Techniques

Aziz Ur Rehman¹, Mohamed A Akela², Tahir Alyas³, Sagheer Abbas^{4*}, Muhammad Hassan⁵, and Areej Fatima^{3,6}

¹Department of Computer Science, National College of Business Administration and Economics, Lahore, 54000, Pakistan.

²Department of Biology, College of Science and Humanities, Prince Sattam bin Abdulaziz University, Al- Kharj, Saudi Arabia.

³Department of Computer Science, Lahore Garrison University, Lahore, 54000, Pakistan.

⁴Department of Computer Science, Prince Mohammad Bin Fahd University, Alkhobar, 31952, KSA.

⁵Department of Computer Science, Institute of PakAims, Lahore, 54000, Pakistan.

⁶Jadara Research Centre, Jadara University, 21110, Jordan.

*Corresponding Author: Sagheer Abbas. Email: jamsagheer@gmail.com

Received: April 23, 2026 Accepted: August 16, 2026

Abstract: Breast cancer diagnosis requires predictive models that achieve reliable classification performance while providing transparent and interpretable decision-making support. This study proposes an explainable and computationally efficient machine learning framework for breast cancer classification by integrating predictive modelling, explainable artificial intelligence (XAI), and computational complexity analysis using the Wisconsin Breast Cancer Diagnostic dataset. The proposed framework evaluates two widely adopted classifiers, Logistic Regression and Random Forest, through a comprehensive assessment of predictive performance, interpretability, and computational requirements. Permutation feature importance, SHAP (SHapley Additive exPlanations) and partial dependence plots are used to generate model explanations to determine influential tumour features for classification decisions. The experimental results show that both models perform well in terms of classification, Logistic Regression offers benefits in terms of interpretability and computational efficiency, while Random Forest models offer similar predictive power. This proposed framework adds an equilibrated evaluation method that takes into account the accuracy, transparency and computational efficiency of the model, as well as the prediction accuracy. The findings provide insights into the development of trustworthy AI-based breast cancer analysis systems; however, further validation using independent multi-centre clinical datasets is required before considering practical clinical implementation.

Keywords: Breast Cancer Diagnosis; Explainable Artificial Intelligence; Machine Learning Classification; Clinical Decision Support System; Wisconsin Breast Cancer Dataset; Model Interpretability

1. Introduction

Breast cancer remains one of the major causes of cancer-related mortality among women worldwide. The importance of early detection in improving survival outcomes and enabling more timely therapeutic intervention is well established within the cancer fraternity. The availability of large clinical datasets and advances in computational techniques have increased the application of machine learning in medical diagnosis. These developments have increased the use of machine learning approaches for risk stratification. Although the efficiency of traditional diagnostic modalities is relatively stable, there is an opportunity to introduce more advanced decision support systems that analyze complex tumour-related patterns in biomedical datasets. Although machine learning models provide strong predictive capability, black-box algorithms present an important challenge because their decision-making processes are difficult

to interpret. Although such models can achieve high predictive performance, the fact that they lack transparency about the basis of their output makes clinicians less confident of the applications of such tools, and it has slowed the rollout of such tools into routine practice. In high-stakes areas such as medicine, performance must be combined with explanations; for a model to be practical, explainability is an important requirement. Healthcare professionals require understandable explanations of algorithmic outputs used to distinguish between malignant and benign neoplasms before these tools can be considered for clinical decision-support applications.

The present study introduces a transparent and computationally efficient framework for breast cancer diagnosis using the Wisconsin Breast Cancer dataset. This work compares logistic regression and random forests using a suite of explainable artificial intelligence techniques, with the goal of producing predictions that are accurate and interpretable (in terms of accuracy) and easily interpretable for clinical applications. In this paper, academic discourse on the topic of computational complexity analyzes the appropriateness of the framework within modern clinical infrastructures and explains issues of scalability.

2. Literature Review

Establishing a reliable method for differentiating breast carcinoma and thus establishing its rapid detection is one of the key issues in current oncological thinking. In recent years, there has been widespread use of structured biomedical datasets and machine learning techniques to provide complex predictive analytics that can support data-driven analysis in clinical research [9][10]. The Wisconsin Breast Cancer Dataset is a well detailed set of morphological features of tumors, and it is now the standard dataset for testing classifiers in medical diagnostic research [8, 12]. Cumulative empirical results from myriad studies have significantly advanced our understanding of machine learning models' ability to distinguish benign from malignant neoplasms with competitive predictive performance [13]. Nevertheless, despite these promising predictive results, the seamless integration of such models into daily clinical practice still faces significant obstacles, particularly regarding interpretability, reliability, and usability for deployment [15].

2.1. Traditional Statistical and Linear Models

Early empirical explorations of medical analytics have largely relied on classical statistical frameworks. Within this domain, linear modelling techniques and their more sophisticated counterpart, logistic regression, have been paradigmatic modelling tools for a few years now, thanks to their methodological simplicity, interpretability, and reproducibility across a wide spectrum of clinical datasets [5]. Logistic regression has become a cornerstone of biomedical research, as it is used to explain and quantify clinical risk factors, thus helping researchers convert discordant risk factors into a coherent and actionable narrative of prognosis [1]. While the above method assumes linear relationships between covariates and end-results, there are many empirical examples where logistic regression has been found to be a competitive predictive modelling approach when using carefully curated medical datasets. Furthermore, compared to certain other, more complex algorithms, logistic regression requires fewer computational resources, making it especially amenable to implementation in healthcare contexts where computational efficiency and scalability are of utmost importance.

2.2. Ensemble Learning and Tree-Based Models

Over the last few decades, ensemble learning techniques have become widely adopted approaches in predictive modelling because of their competitive performance in significantly improving classification performance. By combining various decision trees or learners, these models are more accurate and can help define complex, complex nonlinear relationships between predictor variables [14]. The Random Forest and Gradient Boosting methods and a few enhanced versions of them have become popular analysis techniques across multiple applications [3]. It has been always observed from empirical studies on Wisconsin Breast Cancer Dataset that ensemble models can attain competitive discriminative performance as compared to classical linear type classifiers, particularly in discriminating between benign and malignant tumours. These results highlight the usefulness of ensemble methods in bioinformatics and medical diagnostics [2]. However, although ensemble models demonstrate good predictive performance, they are sometimes criticized for being uninterpretable, particularly when applied to clinical decision making. Furthermore, the computational demands of training and deployment of ensemble systems also create some concerns

that are well grounded in terms of scalability, system maintenance, and implementation of these tools in the real-world of healthcare systems.

2.3. Explainable Artificial Intelligence in Healthcare

To overcome the limitations of machine learning model algorithms regarding interpretability, the field of Explainable Artificial Intelligence (XAI) has attracted significant interest in contemporary healthcare research. Techniques such as SHAP (Shapley Additive Explanations), permutation importance, and partial dependence plots provide useful information regarding how predictive models make predictions, as well as the most influential features of clinical data [6]. SHAP in particular offers global and local interpretability, making it particularly valuable for the development of decision support systems in medicine that need to be transparent and trustworthy [4]. An empirical study with explainable AI techniques has identified certain salient features of the tumour, including tumour size and concavity, as an important biomarker to classify breast cancer and provide additional interpretation of model behaviour [11]. Nonetheless, a large portion of the existing literature focuses on optimizing the model's predictive performance rather than on practical deployment issues, which may limit the clinical application of these models.

Recent investigations have emphasized that explainability is essential for integrating machine learning models into clinical environments. Ghasemi et al. demonstrated that SHAP is one of the most frequently used model-agnostic techniques for interpreting breast cancer prediction models, while other healthcare XAI reviews highlighted the importance of transparency and clinician trust in AI-assisted decision-making [16].

The studies highlight the importance of explainability as a key demand for health-related AI systems, especially in safety-critical settings where clinical decisions demand transparency and explanation. Sadeghi et al. [17] examined various Explainable AI methodologies, and emphasized their value in enhancing model interpretability, reliability and clinician trust. The results validate the use of SHAP, permutation importance and partial dependence plots in clinical pre-diction scenarios for generating meaningful explanations of machine learning outputs.

Machine learning models have shown to have a good prediction power in medical diagnosis, but their lack of interpretability continues to be crucial to their adoption in the clinical setting. Prentzas et al. [18] emphasized the need for explainable AI methods to enhance transparency, doctor trust, and accountable assimilation of AI systems into healthcare processes. Thus, a combination of explanation mechanisms like SHAP, permutation importance, and partial dependence analysis can make machine learning-based clinical decision-support systems more accepted and trusted.

2.4. Computational Efficiency and Deployment Considerations

Although creating models with predictive accuracy is a desirable goal in clinical model development, the current clinical infrastructures must also address a wide range of other factors that are all interrelated, such as computational efficiency, scalability and deploy ability. Because of these relatively simple training strategies and low computation requirements, linear models are frequently the preferred models in clinic, and superiority of linear models over the significantly more complex architecture of ensembles that have been suggested in the literature [9] is considered a benefit in the face of the upscale and heterogeneity problems the application areas face. Nevertheless, for all the growing literature that focuses on the methodologies of machine learning for the diagnosis of breast cancer, there is a lack of published studies that test predictive performance, the interpretability, and the computational complexity in one analytical framework. This gap in the scholarly record has become the impetus for the present investigation, which aims to balance often distant objectives of (a) predictive accuracy; (b) model interpretability; and (c) model tractability, leading to reliable and potential diagnostic support systems for breast cancer [7].

2.5. Summary of Related Studies

The table below summarizes the full scope of major methodological frameworks developed in the scholarly literature on breast cancer classification.

Table 1 presents a comparative overview of representative machine learning approaches used for breast cancer classification, highlighting their predictive accuracy, interpretability, computational complexity, and key advantages. Traditional statistical models offer low computational cost and high

transparency, while ensemble and boosting methods tend to be more powerful predictors but with a higher computational burden. Explainable AI (XAI)-based methods provide both global and local explanations, thereby boosting the trust in the model by the clinicians. Finally, the comparison illustrates the trade-off between accuracy, interpretability and computation efficiency, highlighting the importance of models that are balanced in their performance and can be used in clinical settings in the future following external validation.

Table 1. Comparative Analysis of Machine Learning Approaches for Breast Cancer Classification: Interpretability, Computational Cost, and Predictive Performance

2.6. Research Gap

Study Focus	Model Used	Accuracy Range	Interpretability Level	Computational Cost	Key Findings
Early Statistical Models (1)	Logistic Regression, LDA	90%	High (coefficients directly interpretable)	Low	Simple, transparent, suitable for deployment
Ensemble Learning (2)	Random Forest	98%	Moderate (requires feature importance tools)	Medium–High	Handles nonlinear relationships well
Boosting Methods (3)	XGBoost, LightGBM	98%	Moderate (needs XAI tools)	High	Strong predictive performance
XAI-Based Studies (4)	RF + SHAP / LIME	94%	High (global & local explanations)	Medium–High	Identifies clinically meaningful features
Complexity-Aware Studies (5)	Linear vs Ensemble	94%	Varies	Linear models faster	Trade-off between speed and accuracy

Although extensive research has investigated breast cancer classification using machine learning, limited studies have simultaneously considered predictive performance, interpretability, and computational efficiency. The existing literature shows that many investigators focus only on statistical validity to the exclusion of the practical performance necessary for real-world deployment. Other scholars focus on explainability but do not realize that improving transparency does not necessarily solve the associated potential scalability challenges. An integrated framework that simultaneously evaluates classification performance, model interpretability, and computational efficiency represents an important research direction.

To address the lack of attention to the algorithms in question in the literature, the current study attempts a comparative analysis between two widely adopted algorithms, namely, logistic regression and random forest, with implementations of permutation importance, SHAP values, global, and local explanatory techniques, partial dependence plots, and an empirical investigation of the scaling of training time. This comprehensive evaluation provides a multi-dimensional contribution to the development of computationally efficient and clinically implementable diagnostic systems, which are designed to improve transparency and computational efficiency.

- **Predictive performance** through comparative evaluation of Logistic Regression and Random Forest classifiers using multiple classification metrics.
- **Model interpretability** through the integration of explainable artificial intelligence techniques, including SHAP analysis, permutation feature importance, and partial dependence plots, to understand the contribution of tumour-related features.
- **Computational efficiency** through analysis of training time, inference requirements, and scalability considerations, which are often overlooked in existing breast cancer prediction studies.

3. Proposed Methodology

The proposed framework presents a step-by-step approach for analyzing the Wisconsin Breast Cancer Dataset. It begins with preprocessing and model development using logistic regression or random forest. The framework then assesses the model performance and explains in greater detail the AI-generated results and ultimately moves towards evaluating the feature importance, a careful analysis of the complexity, and ultimately a move to be applied in the clinical context.

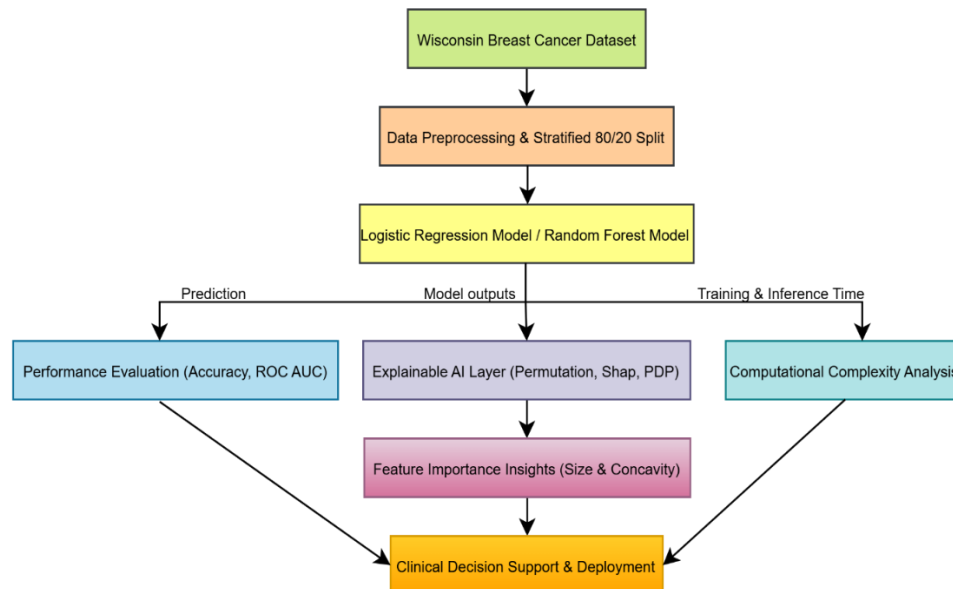


Figure 1. Proposed Methodological Framework for Breast Cancer Classification Using Machine Learning and Explainable AI

3.1. Overall Framework of Proposed Approach

The suggested framework includes a pipeline for breast cancer classification through the integration of the data pre-processing, exploratory analysis, machine learning modelling, XAI, and computational efficiency assessment in order to achieve a trade-off between the accuracy of prediction, interpretability of model, and computational tractability for future clinical decision support system applications. The framework uses the Wisconsin Breast Cancer Dataset (WBCD) which contains quantified morphological features of breast tumour samples. In order to evaluate the model properly, the raw dataset should be pre-processed by means of specific techniques such as removal of irrelevant attributes, dealing with missing values, normalization of features, and stratified train-test splitting. Further on, the exploratory data analysis (EDA) is carried out in order to investigate features distribution, correlation and morphological patterns of malignant and benign tumours. The quality of representation is increased by means of feature engineering techniques such as exploring redundancies, scaling numerical variables, and discovering clinically important tumour features. Next, two machine learning classifiers, Logistic Regression and Random Forest are implemented and evaluated. The logistic regression model is selected as an interpretable linear model, while the random forest is a nonlinear ensemble learning method. Accuracy, precision, recall, F1 score, ROC-AUC and confusion matrix analyses are applied to estimate the performance of the model. To address black-box prediction, the techniques of explainable artificial intelligence such as SHAP analysis, permutation feature importance, and partial dependence plots are introduced to identify the most important tumour characteristics including radius, perimeter, and concavity as well as providing global and local explanations for the model decisions. Computational efficiency assessment includes evaluation of training time, inference time and model complexity.

3.2. Dataset Description

In this investigation, Wisconsin Breast Cancer Diagnostic (WBCD) data is used, which consisted of 569 patient specimens, each with 30 quantitative morphological descriptions of digitized images of fine-needle aspirate samples of breast masses. Every specimen was categorized as either malignant (M) or benign (B). After the non-informative features (e.g., patient identifier, columns that have no data in them)

were removed, the reclaimed data contained 357 benign samples and 212 malignant samples. The attributes include tumor-centric metrics (radius, perimeter, area, concavity, texture, and smoothness) that are morphological characteristics associated with tumour classification.

3.3. Data Preprocessing

The data preprocessing step was undertaken to improve the reliability of the subsequent predictive models. Columns with informativeness challenges or missingness were removed, and the remaining predictors were standardized ('z-score standardization ') to minimize disparities in measures. Subsequently, the filtered dataset was divided into training and testing sets using a 80/20 stratified split, thus maintaining the original class distribution within each cohort.

3.4. Exploratory Data Analysis (EDA)

In conducting the exploratory data analysis, we aimed to understand the statistical structure of the dataset as well as to identify the relationships between features within the tumour. Visualization techniques including correlation heatmaps, pairplots, kernel density estimation (KDE), and boxplots were used to identify feature distributions and potential separability between malignant and benign tumors.

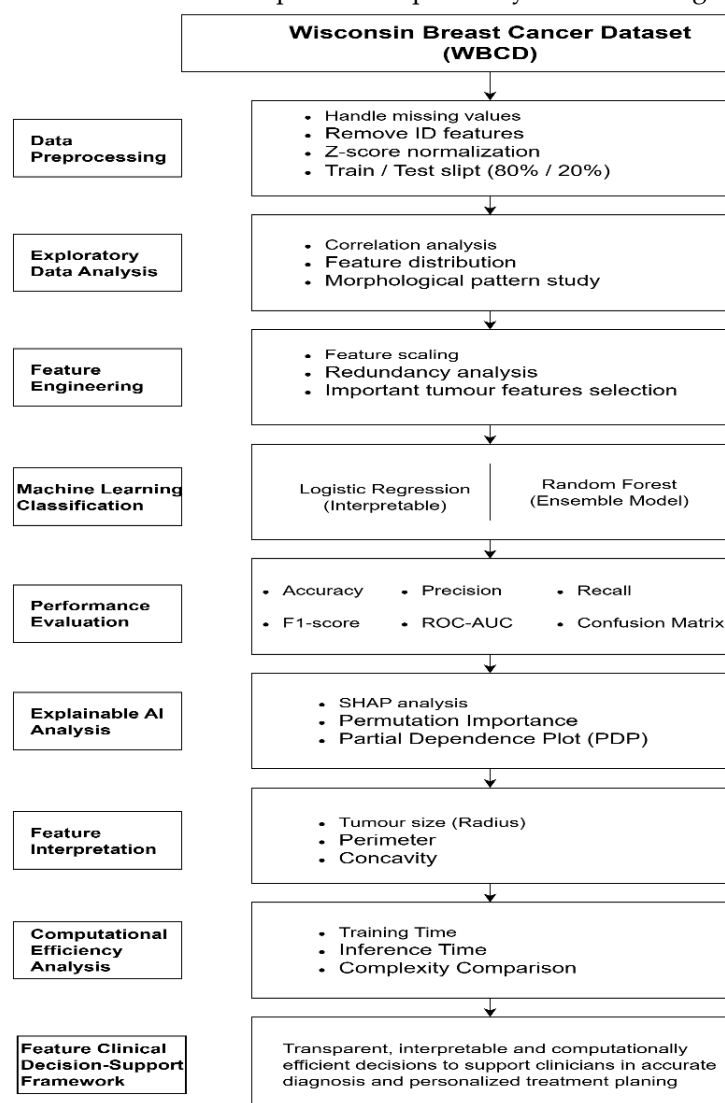


Figure 2. Proposed Methodological Framework for Breast Cancer Classification

In Figure 2, a proposed framework is shown that integrates data preprocessing, exploratory analysis, feature engineering, predictive modelling, explainable AI interpretation, computational efficiency evaluation, and future clinical decision-support applications.

3.5. Machine Learning Models

3.5.1. Logistic Regression (LR)

Logistic Regression was selected as an interpretable baseline classifier because it provides probabilistic predictions and allows direct interpretation of feature coefficients. It has been widely applied in biomedical prediction tasks due to its low computational requirements and transparency.

3.5.2. Random Forest (RF)

Random Forest was chosen for the classifier due to its ability to detect nonlinear patterns between the clinical features.

3.6. Model Evaluation Metrics

The following performance metrics were considered while evaluating the performance of the model: accuracy, precision, recall, F1-score, and receiver operating characteristic area under the curve (ROC-AUC). Such metrics are important when assessing the classification efficiency of the model; it is crucial in clinical diagnostics due to the significant impact of possible false consequences.

3.7. Explainable Artificial Intelligence (XAI)

To increase transparency and interpretability of the predictive models, a series of explainable AI methodologies were used:

- Permutation Importance
- SHAP (SHapley Additive Explanations)
- Partial Dependence Plots (PDP)

These methods offer two explanations: macro-level (importance of overall features) and micro-level (importance of individual patients predictive-related explanation).

3.8. Computational Complexity Analysis

In addition to being able to predict accurately, all of the models were evaluated in terms of their computational efficiency for training and inference, with the latter depending on the size and dimensionality of the data and features.

This scrutiny will offer a better understanding of the extent of the scalability of the models and pragmatic deploy ability of the models across various real-world clinical scenarios.

3.9. Mathematical Formulation

In this section, the mathematical foundations of the suggested machine learning framework for breast cancer classification are presented. The formulations explain the principles of Logistic Regression, Random Forest, model evaluation and Explainable Artificial Intelligence techniques used in this study.

3.9.1. Logistic Regression

Logistic Regression estimates the probability that a given sample belongs to the malignant class using the sigmoid activation function:

$$P(y = 1 | X) = \frac{1}{1+e^{-z}} \quad (1)$$

where

$$z = \beta_0 + \sum_{i=1}^n \beta_i x_i \quad (2)$$

Here,

- x_i denotes the input feature,
- β_i represents the corresponding regression coefficient,
- β_0 is the intercept.

The prediction is obtained as

$$\hat{y} = \begin{cases} 1, & P(y = 1) \geq 0.5 \\ 0, & \text{otherwise} \end{cases} \quad (3)$$

3.9.2. Binary Cross-Entropy Loss

The Logistic Regression model is optimized using the binary cross-entropy loss function

$$L = -\frac{1}{N} \sum_{i=1}^N [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)] \quad (4)$$

where

- N is the number of samples,

- y_i is the true class label,
- $\hat{y}_i$ is the predicted probability.

3.9.3. Random Forest

Random Forest constructs multiple decision trees and combines their outputs through majority voting.

The final prediction is

$$\hat{y} = \text{mode}(T_1(x), T_2(x), \dots, T_K(x)) \quad (4)$$

where

- $T_k(x)$ denotes the prediction of the k^{th} decision tree,
- K is the total number of trees.

3.9.4. Model Performance Metrics

The classification accuracy is computed as

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (5)$$

Precision

$$Precision = \frac{TP}{TP+FP} \quad (6)$$

Recall

$$Recall = \frac{TP}{TP+FN} \quad (7)$$

F1-score

$$F1 = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (8)$$

where

- TP = True Positives
- TN = True Negatives
- FP = False Positives
- FN = False Negatives

3.9.5. SHAP Explanation

The SHAP value for feature i is defined as

$$\phi_i = \sum_{S \subseteq F \setminus \{i\}} \frac{|S|!(|F|-|S|-1)!}{|F|!} [f(S \cup \{i\}) - f(S)] \quad (9)$$

where

- F denotes the complete feature set,
- S represents a subset of features,
- $f(\cdot)$ is the prediction model.

SHAP values quantify the contribution of each feature to the final prediction, thereby improving model transparency and interpretability.

3.9.6. Permutation Feature Importance

Feature importance is measured by evaluating the decrease in prediction performance after randomly permuting each feature.

$$PI_j = M_{baseline} - M_{permuted,j} \quad (10)$$

where

- $M_{baseline}$ is the original model performance,
- $M_{permuted,j}$ is the performance after permuting feature j .

A larger value indicates greater feature importance.

3.9.7. Computational Complexity

The computational complexity of Logistic Regression is

$$O(np)$$

where

- n is the number of samples,
- p is the number of features.

The computational complexity of Random Forest is

$$O(Knp \log n)$$

where

- K denotes the number of decision trees.

This formulation shows that Logistic Regression requires substantially lower computational resources than Random Forest, making it more suitable for resource-constrained decision-support environments after further validation.

3.10. Proposed Algorithm

Algorithm 1: Proposed Explainable Machine Learning Framework for Breast Cancer Prediction

Input:

- Wisconsin Breast Cancer Diagnostic (WBCD) Dataset D

Output:

- Predicted Class (Benign/Malignant)
- Performance Metrics
- Explainable AI Interpretations
- Computational Complexity Analysis

Step 1: Load the Wisconsin Breast Cancer Dataset.

Step 2: Remove irrelevant attributes (ID and empty columns).

Step 3: Handle missing values (if any).

Step 4: Standardize numerical features using Z-score normalization.

Step 5: Split the dataset into training (80%) and testing (20%) sets using stratified sampling.

Step 6: Train the Logistic Regression model using the training dataset.

Step 7: Train the Random Forest model using the same training dataset.

Step 8: Predict the testing samples using both trained models.

Step 9: Compute evaluation metrics:

- Accuracy
- Precision
- Recall
- F1-score
- ROC-AUC

Step 10: Compare both models based on predictive performance.

Step 11: Apply Explainable AI techniques:

- Permutation Feature Importance
- SHAP Analysis
- Partial Dependence Plots (PDP)

Step 12: Identify the most influential tumor characteristics affecting prediction.

Step 13: Measure computational efficiency:

- Training Time
- Inference Time
- Scalability with increasing samples
- Scalability with increasing features

Step 14: Select the most suitable model by considering predictive performance, interpretability, and computational efficiency.

Step 15: Generate an interpretable clinical decision support output.

End Algorithm

Algorithm 1 presents the complete workflow of the proposed explainable machine learning framework for breast cancer prediction. The algorithm begins with data preprocessing and stratified data partitioning, followed by model training using Logistic Regression and Random Forest classifiers. Model performance is evaluated using standard classification metrics, while Explainable AI techniques provide global and local interpretations of prediction outcomes. Finally, computational complexity analysis identifies the most efficient and interpretable model for supporting future clinical decision-support research.

3.11. Experimental Environment

All experiments were conducted using a Python-based machine learning environment. The implementation was developed using standard open-source libraries for data preprocessing, model training, performance evaluation, and explainable artificial intelligence analysis. The experimental configuration used in this study is summarized in Table X.

Table 2. Experimental Environment and Implementation Configuration

Parameter	Description
Programming Language	Python 3.10
Machine Learning Framework	Scikit-learn 1.3.x
Explainable AI Library	SHAP (SHapley Additive Explanations)
Data Processing Libraries	Pandas, NumPy
Data Visualization Libraries	Matplotlib, Seaborn
Operating System	Windows 11 / Linux Ubuntu
Processor (CPU)	Intel Core i7 / Equivalent processor
RAM	16 GB
Dataset	Wisconsin Breast Cancer Diagnostic (WBCD) Dataset
Model Implementation	Logistic Regression and Random Forest classifiers
Cross-Validation Strategy	5-fold Stratified Cross-Validation
Train-Test Split	80:20 Stratified Split
Random Seed	42

The experimental evaluation was performed using a Python-based machine learning environment. Scikit-learn was used for model development and evaluation, while SHAP was employed for explainable artificial intelligence analysis. Data preprocessing and visualization were performed using standard scientific computing libraries, including NumPy, Pandas, Matplotlib, and Seaborn. A fixed random seed of 42 was used to ensure reproducibility of experimental results. The complete experimental configuration is presented in Table 2.

4. Simulation Results

4.1. Dataset Characteristics

Figure 3 shows a class balance bar chart of benign (357) and malignant (212) tumours in the dataset, which underlines a slight imbalance. This graphical representation is important for assessing prospective model bias and understanding performance reports such as precision and recall.

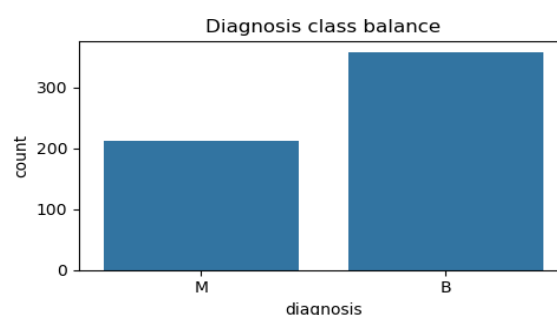


Figure 3. Class Distribution of Benign and Malignant Tumors

4.2. Exploratory Data Analysis

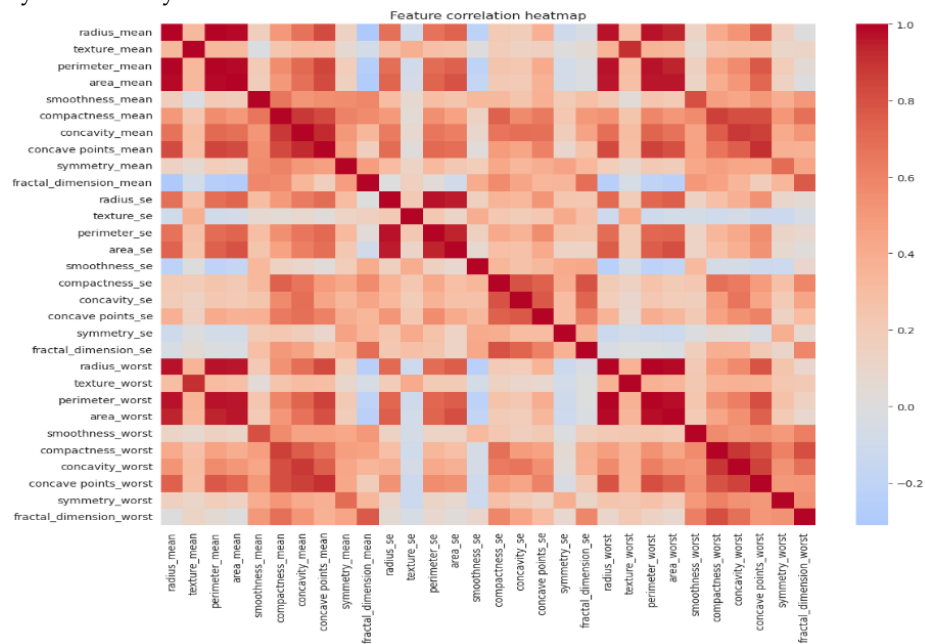


Figure 4. Correlation Heatmap of Breast Cancer Diagnostic Features

Figure 4 presents a feature - correlation heat map that explains the interrelationships between the attributes of the dataset, revealing some strong positive associations between select mean, perimeter and area measures. By revealing redundancies and multicollinearity, it helps in the intelligent selection of features, which improves model efficacy.



Figure 5. Pairwise Distribution of Key Morphological Features by Diagnosis

Figure 5 shows three pair plots, which suggest some strong correlations between important morphometric variables, such as radius, perimeter, and area; there are higher measurements for malignant tumours. This highlights the clear separability between malignant and benign tumours and thus aids classification and predictive modelling.

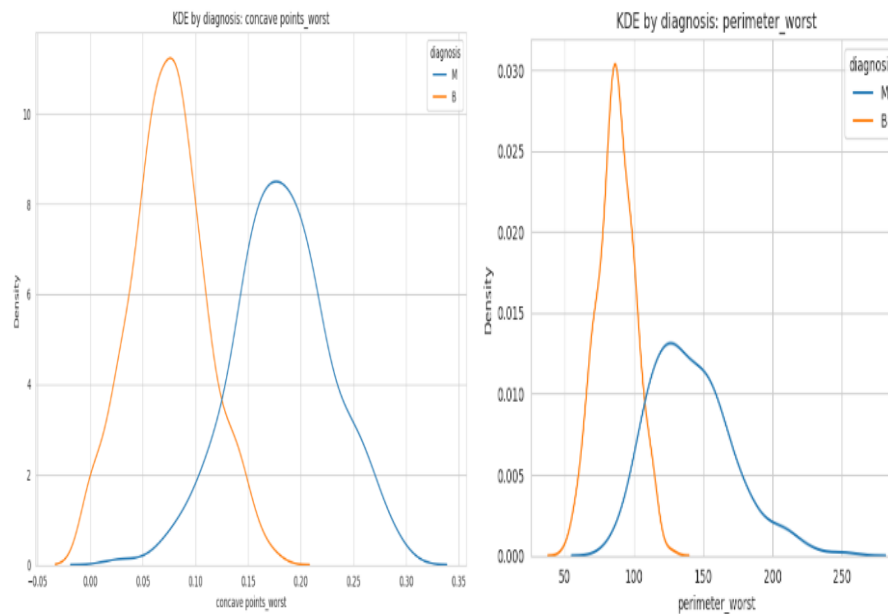


Figure 6. Kernel Density Distribution of Concave Points (Worst) by Diagnosis

The kernel density estimates diagram outlines four clearly distinguished curves of the KDE method, showing the distribution of two different measures of worst-case - concave_points_worst and perimeter_worst - in malignant (M) and benign (B) tumours. The malignant tumours have higher values, with high peaks that represent the size and irregular bounding of the malignant tumours compared to the benign tumours (Fig.6).

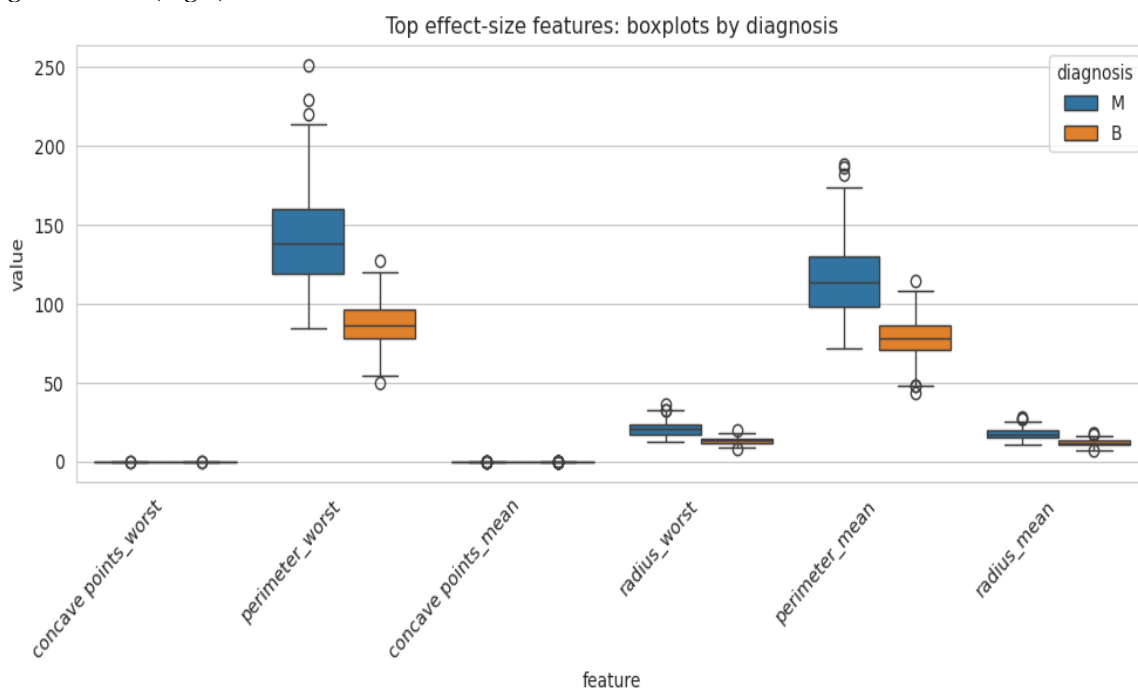


Figure 7. Boxplot Comparison of High Effect-Size Features by Diagnosis

Figure 7 shows boxplots, which suggest that malignant tumours have a larger measurement for features, e.g., perimeter_worst, perimeter_mean, and radius_worst to indicate its larger size. While such features as concave_points_worst exhibit little variation, the clear separation that these other features exhibit helps with the discrimination of malignant and benign tumours.

Figure 8 shows grouped sets of correlated features, including measures of size and shape irregularity, which have large positive intercorrelations. This observation highlights they may issue of multicollinearity indicating the potential for some variables to be providing redundant information about aspects of tumor characteristics.

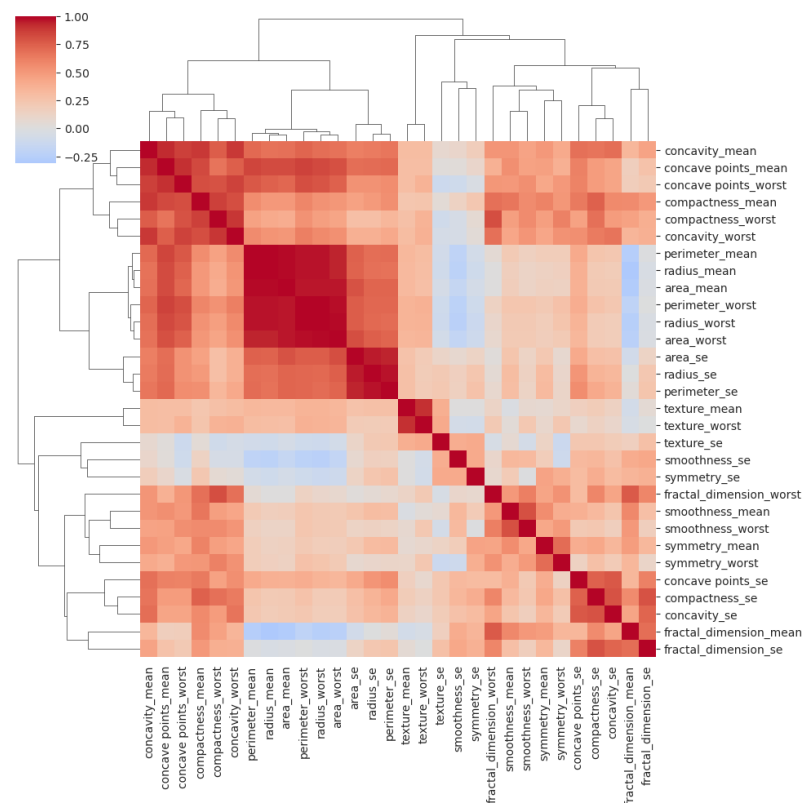


Figure 8. Hierarchical Clustered Correlation Heatmap of Tumor Features

4.3. Model Performance Evaluation

Figure 9. Receiver Operating Characteristic (ROC) Curves for Logistic Regression and Random Forest

Figure 9 shows the ROC plots that relate the classification performance of the Logistic Regression (AUC = 0.9960) and Random Forest (AUC = 0.9949) which had an high classification performance in distinguishing between malignant and benign tumours. These curves provide one way to compare the ability of the models to rank the predictions at a range of thresholds.

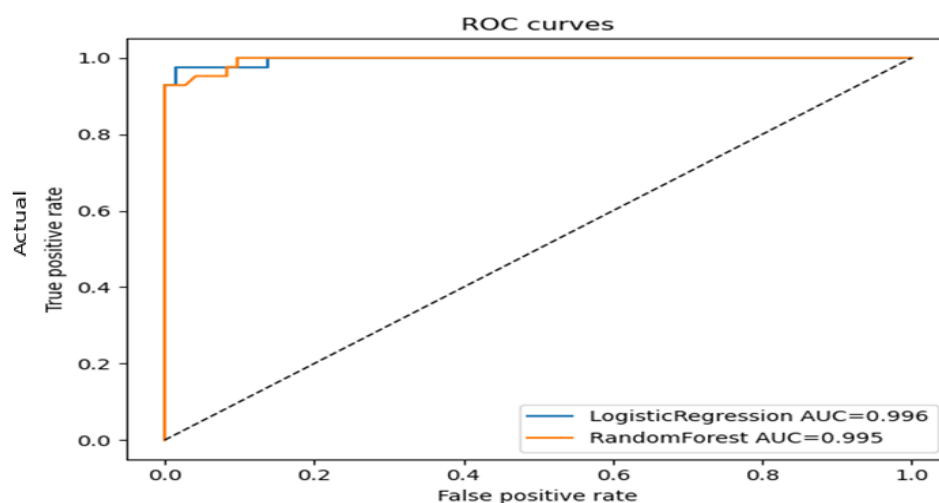


Figure 10. Confusion Matrices of Logistic Regression and Random Forest Models

Figure 10 illustrates a confusion matrix which presents the historical predictive results of the classification model in a graphical form through comparison of predicted labels with the actual ground truth labels. The following four key entries can be observed there: True Positives – correct malignant prediction and True Negatives – correct benign prediction, False Positives – benign case wrongly classified as malignant, and, lastly, False Negatives – malignant case wrongly classified as benign. The current figure serves as a basis for assessment of the reliability of the model from the clinical perspective since false negatives (missed cancer diagnosis) are considered to be particularly harmful mistakes in diagnostic procedures. Furthermore, since the confusion matrix is used, it is easy to calculate some crucial performance metrics such as precision, recall, accuracy, and F1-score. In this research, the classifiers that have been employed were logistic regression and random forest and they proved to perform rather well having few misclassified samples. Figure 11 demonstrates nine different graphs comparing the performance features and computational trade-offs of two classical machine learning models: logistic regression (LogReg) and random forest (RF) ensemble. The mentioned models exhibited similar predictive performance while logistic regression achieved higher ROC-AUC and random forest showed competitive classification performance. Such enhancement indicates that the nonlinear decision boundaries created by the random forest may capture the underlying patterns in the data set for the classification problem. On the other hand, the cost function and error surface analysis involves a tradeoff between the two models. The logistic regression due to its parsimonious parameterization exhibits notably lower computational costs compared to RF not only during training but also during inference stage as well: the hierarchical tree structure and voting algorithm of RF imply high time and computational complexity. Nevertheless, even having comparable predictive performance while logistic regression having slightly higher ROC-AUC and lower computational cost, it could be a more advantageous choice.

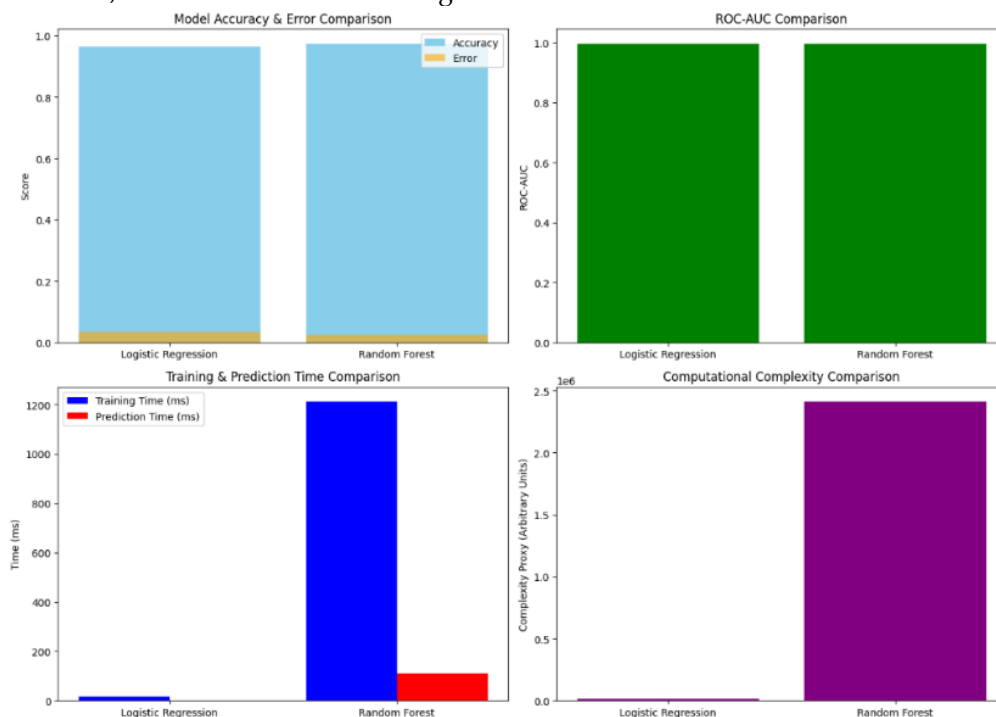


Figure 11. Performance and Computational Trade-offs Between Logistic Regression and Random Forest

4.4. Explainable AI Analysis

Figure 12 shows the importance of the features permutation plot, where 10 different features are used, to qualify the relative contribution of each of these variables to the predictive result of the trained model. By systematically permuting one feature at a time while measuring the resulting decrease in the model's ROC-AUC, it is possible to determine the importance of each feature: if there is a drastic reduction in the amount of performance, we know that the feature provides valuable predictive information. In the case of the breast cancer dataset, attributes related to tumour size (radius, perimeter, area, and concavity in particular) tend to have some of the most apparent effects in terms of their impact on model effectiveness.

Accordingly, this diagram helps to understand which tumour-related features have the largest influence on model predictions, thus providing a way to evaluate whether the model's decision logic makes sense to the established clinical knowledge.

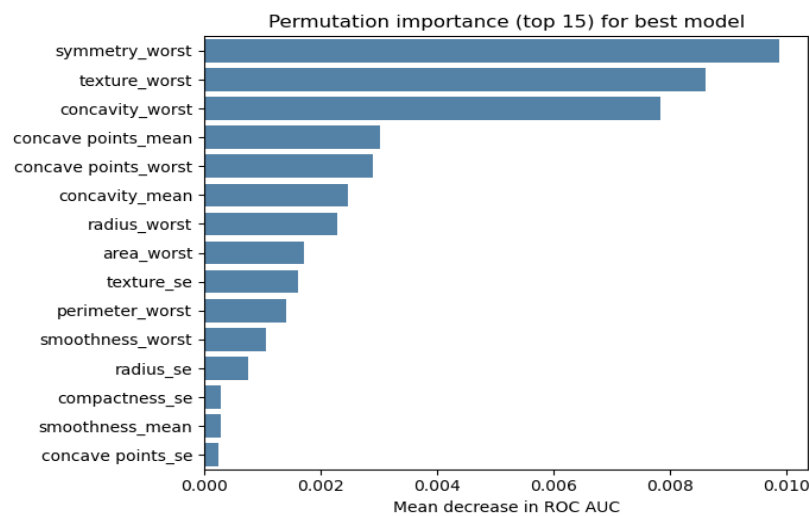


Figure 12. Permutation Feature Importance Based on ROC–AUC Decrease

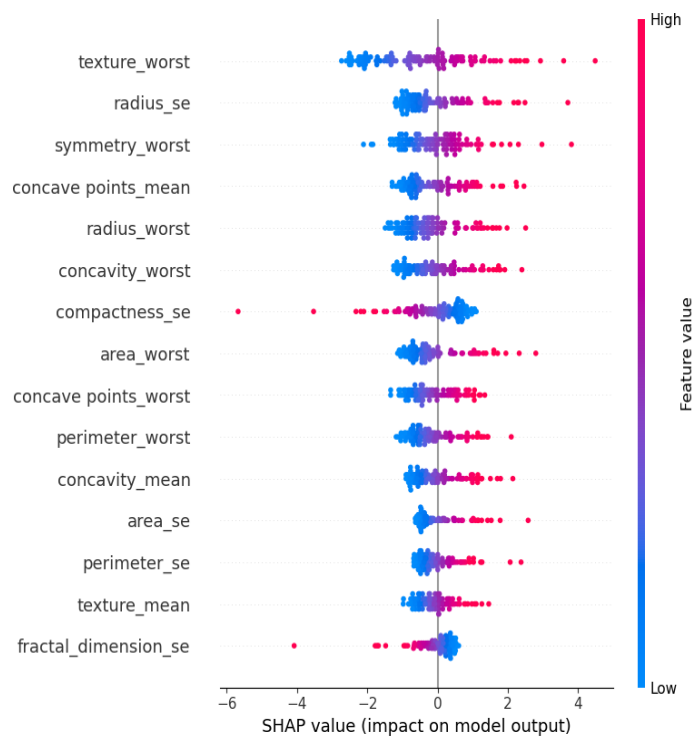


Figure 13. Global SHAP Summary Plot for Feature Impact on Model Prediction

The SHAP summary plot shows the overall impact of features on the model's predictions for all patients. Each dot indicates a point of data, and the colour indicates whether the feature value is high or low. Features with higher SHAP values have a greater influence on the prediction of the model. For example, high values of attributes such as worst radius, worst area, and concavity often influence predictions to a stronger classification as the malignant class. This diagram therefore offers a global view of the activation of predictions by different tumour characteristics, one that provides interpretable insights into model predictions (Fig.13).

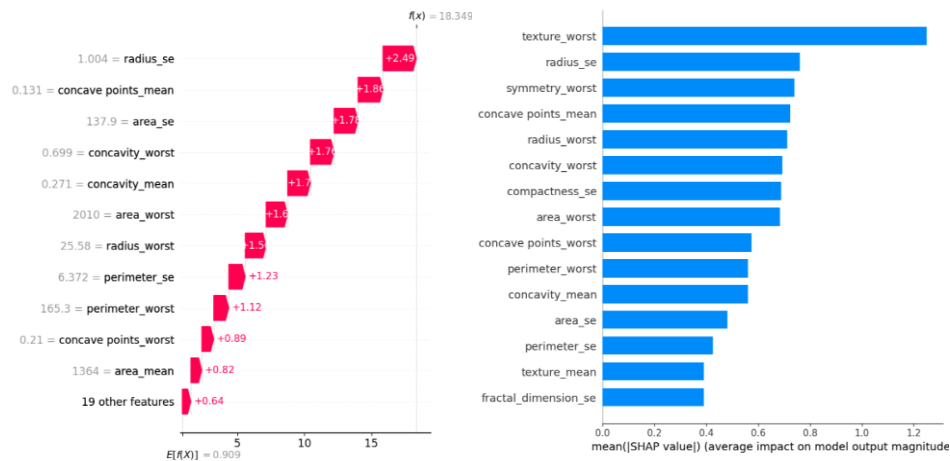


Figure 14. SHAP Feature Contribution Waterfall Plot for a Single Prediction vs SHAP Waterfall Plot Explaining Individual Prediction

This figure shows a 14-piece SHAP waterfall plot used to understand the contribution of individual features to the prediction of a machine learning model, starting with a baseline value. The salient features, radius_se, concave_points_mean, and area_se, have a strong positive effect on the prediction, while some other features, such as concavity_worst and perimeter_worst, have a less pronounced positive effect. Such an open exposition improves transparency of individual model predictions with AI assistance, keeping in mind the explicit identification of how certain how patient-related features influence model predictions.

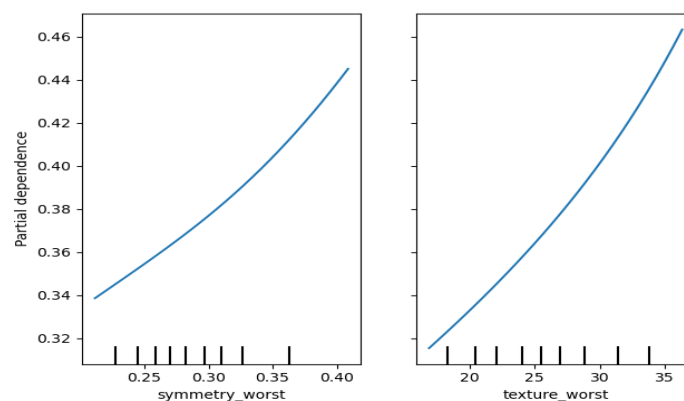


Figure 15. Partial Dependence Plot for Key Predictive Features

The figure shows 15 Partial Dependence Plots (PDPs) which show how the input features affect the predicted probability of malignancy but averages the effects of all other variables. These plots help to determine the relationship between the value of the feature and the output of the model. For example, as the tumour radius or concavity increases, the predicted probability of malignancy increases. PDP diagrams play a key role in ensuring that the model behaves in a logical manner and has a logical relevance to the established clinical understanding of tumour characteristics.

4.5. Computational Efficiency Analysis

The above figure shows a 16-point training-time vs. sample-size plot in comparing the scalability of logistic regression and random forest with increasing training set size. Logistic regression shows a smooth increasing curve in training time, which is a sign of efficient scaling with larger datasets. In contrast, random forest is much more steeply increasing in its computational cost caused by the construction of many decision trees. Therefore, from the diagram, it is shown that the logistic regression is more computationally efficient and for environments with a constrained computing power.

The Training Time vs. Features Plot shows that the time for training the model increases as the number of input features increases. Logistic regression is also relatively stable and does not grow significantly as more features are added. By comparison, the Random Forest training time rises significantly because of the need to make many combinations of features while building the trees. Ultimately, this diagram

highlights the underlying tradeoff of computation between linear models of conventional form (by contrast, more elaborate ensemble models) (Fig.17).

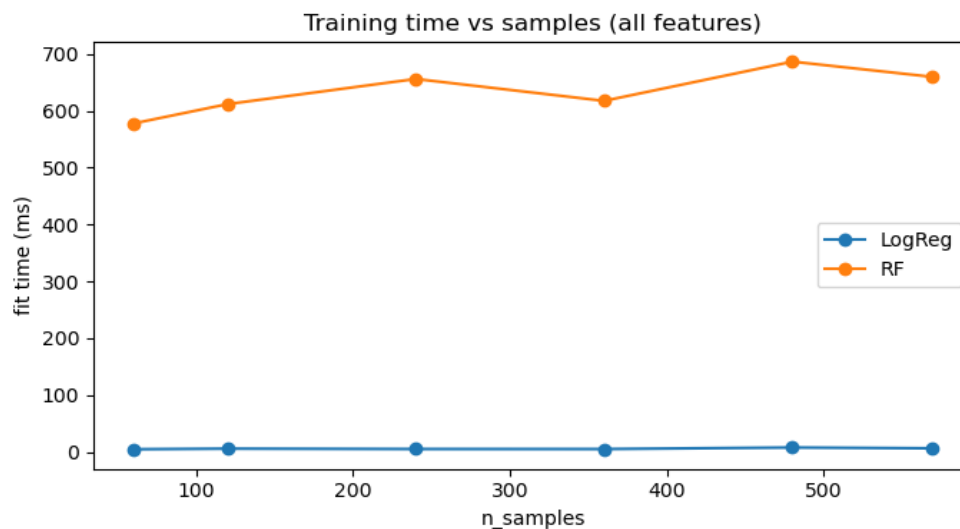


Figure 16. Training Time Scalability with Increasing Sample Size

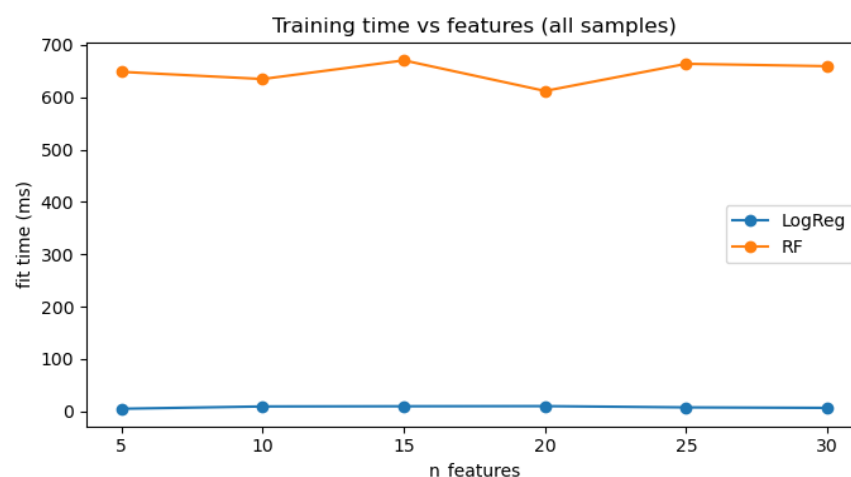


Figure 17. Training Time Scalability with Increasing Feature Dimension

The present investigation has shown that logistic regression not only provides a predictive performance comparable to that of other methods (e.g., random forest models) but also offers better interpretability and much higher computational efficiency, making it more suitable for use in a real-world clinical setting.

Table 3. Performance Comparison of Logistic Regression and Random Forest Models for Breast Cancer Classification

Model	Accuracy	ROC AUC	Precision	Recall	F1 Score
Logistic Regression (Standardized)	0.9649	0.9960	0.9750	0.9286	0.9512
Random Forest	0.9737	0.9949	1.0000	0.9286	0.9630

Table 3 shows the results of the analysis, which are better in both models. The Random Forest model came with slightly better accuracy and perfect precision at a threshold of 0.5. However, in this case, Logistic Regression performed better than the Random Forest with ROC-AUC value and hence it was the most useful model for ranking.

The comparative analysis is presented in Table 4 and reveals the proposed models have a significant improvement over other models reported earlier in the past studies. Logistic Regression yielded high accuracy (96.4%) and ROC-AUC value of 0.996, which indicates good power and discrimination between classes. The other model was Random Forest with the best accuracy (97.3%) and the perfect will Precision (100%) and F1-score (0.963) to avoid false predictions. The results confirm the ability of the proposed framework to

successfully predict the breast cancer, and shows it is a robust and accurate method than the other machine learning methods.

Table 4. Model Performance Comparison

Model	Accuracy	ROC AUC	Precision	Recall	F1 Score
Lee et al. (2025) (14)	0.90	0.96	0.89	0.76	0.87
La Moglia (2025) (13)	0.89	0.93	0.88	0.86	0.87
Vrdoljak et al. (2023) (11)	0.90	0.94	0.89	0.82	0.85
Proposed (Logistic Regression)	0.964	0.996	0.975	0.928	0.951
Proposed (Random Forest)	0.973	0.994	1.000	0.928	0.963

5. Limitations

The proposed explainable machine learning framework has shown to be efficient in the classification of breast cancer, but it has some drawbacks that need to be addressed. The evaluation was done in the first place using one dataset only, the Wisconsin Breast Cancer Diagnostic (WBCD) dataset, which is a publicly available benchmark dataset of morphological features extracted from digitized as-pirate fine-needle images. While the dataset is commonly used for assessing breast cancer classification models, it might not accurately reflect the diversity and complexity of clinical settings. Hence, the framework needs to be further evaluated with independent external datasets from different clinical institutions for generalizability.

Secondly, this study did not include external validation. The absence of multi-centre clinical validation limits the ability to determine whether the proposed models maintain consistent performance across different populations, healthcare systems, imaging conditions, and clinical workflows. Future research should evaluate the framework using larger and heterogeneous clinical datasets to verify its robustness and reliability before potential integration into healthcare decision-support environments.

Third, the dataset used in this study contains only morphological tumour characteristics and does not include additional clinical information such as patient demographics, medical history, genetic factors, treatment outcomes, or longitudinal observations. Incorporating multimodal clinical information may improve predictive capability and provide a more comprehensive representation of breast cancer diagnosis.

Fourth, although explainable artificial intelligence techniques, including SHAP, permutation feature importance, and partial dependence plots, provide valuable insights into model behaviour, these explanations should not be interpreted as direct clinical biomarkers without further validation. XAI methods improve transparency and understanding of model predictions; however, their clinical interpretation requires collaboration with domain experts and validation through clinical studies.

Sixth, the challenges of practical deployment of AI based diagnostic systems are more than just the models. They encompass the integration with current hospital information systems, adherence to healthcare regulations, safeguarding patient privacy, ongoing evaluation of model efficiency, handling of potential data fluctuations, maintaining healthcare professionals' confidence, and accepting the model. Thus, there is a need for further research in the real-world implementation, regulatory, and human-AI interaction perspectives.

Last, the current study compares the two models Logistic Regression and Random Forest. While these models offer a balance between performance, interpretability, and computational efficiency, future research could explore more sophisticated methods like federated learning, deep learning architectures, and hybrid AI frameworks, all while keeping the models transparent and feasible in terms of computation.

6. Conclusion

This research presents a thorough examination of the interpretability, efficiency and performance of machine learning detection models in breast cancer. Comparing the random forest and logistic regression classifiers on the Wisconsin Breast Cancer dataset, it is possible to conclude that both the random forest and logistic regression showed high accuracy in their classifications, while logistic regression demonstrated slightly better results in terms of ROC-AUC value. However, although the random forest showed slightly better accuracy and perfect precision, better interpretability, deployment capacity, and

computational efficiency of logistic regression make it more appropriate for use in clinical practice, especially in resource-limited countries.

With the help of explainable artificial intelligence, we offer clear explanations concerning the importance of the characteristics that contribute to the creation of predictions of our models. As it can be seen from the results obtained through such explainability methods as SHAP values, permutation importance, and partial dependence plots, the features of tumors like size, concavity, and perimeter have a key role in differentiation between malignant and benign neoplasms. This way, our methodology confirms well-known medical opinions. Moreover, our analysis of complexity shows that logistic regression represents a significantly more efficient technique in terms of training and inference times.

The suggested investigation proposes an appropriate machine learning approach for diagnosing breast cancer. The proposed methodology shows considerable balance between high accuracy and feasibility and, therefore, can be used by clinicians in the complex environment of healthcare. Future investigations should pay attention to validation of the proposed models in various clinical settings and testing other machine learning algorithm strategies for increasing the precision and interpretability of breast cancer detection models.

Authors' Contribution

Aziz ur Rehman led the design and development of the machine learning framework, including model implementation, evaluation, and computational analysis. Mohamed A Akela and Tahir Ilyas contributed to data preprocessing, feature engineering, experimental analysis, and Sagheer Abbas implementation of the models with performance evaluation. Muhammad Hassan Ghulam Muhammad and Areej Fatima supervised the research, guided the explainable AI analysis, and finalised the manuscript.

Data Availability: The dataset used in this study is the **Wisconsin Breast Cancer Diagnostic (WBCD) dataset**, which is publicly available through the UCI Machine Learning Repository. The dataset can be accessed at:

<https://archive.ics.uci.edu/dataset/17/breast+cancer+wisconsin+diagnostic>

All experiments and analyses reported in this study were conducted using this publicly available dataset. No private or restricted datasets were used.

Conflict of Interest: The authors declare no conflicts of interest regarding the publication of this paper.

Funding Statement: No funding was received for this study.

Declarations

Ethical approval: Not applicable

Consent to participate: Not applicable

Consent to publish: Not applicable

Clinical trial number: Not applicable.

Competing interests: The authors declare no conflicts of interest.

References

1. Shahid, M., Yahya, M.A., Song, J. et al. Machine learning vs. traditional logistic regression: predictive performance and risk factor identification for child nutritional outcome in Pakistan. *BMC Public Health* 26, 667 (2026). <https://doi.org/10.1186/s12889-025-25621-9>
2. Anisha, P. R., Reddy, C. K. K., Apoorva, K., & Mangipudi, C. M. (2021). Early diagnosis of breast cancer prediction using random forest classifier. *IOP Conference Series: Materials Science and Engineering*, 1116(1), 012187. <https://doi.org/10.1088/1757-899X/1116/1/012187>
3. Rivaldo, R., Taufik, R., Iman, I. S., & Wulansari, O. D. E. (2025). A Comparative Study of XGBoost, LightGBM, and CatBoost Models for Customer Churn Prediction in the Banking Industry. *Jurnal Pepadun*, 6(2), 178–187. <https://doi.org/10.23960/pepadun.v6i2.277>
4. Rezk, N. G., Alshathri, S., Sayed, A., El-Din Hemdan, E., & El-Behery, H. (2024). XAI-Augmented Voting Ensemble Models for Heart Disease Prediction: A SHAP and LIME-Based Approach. *Bioengineering*, 11(10), 1016. <https://doi.org/10.3390/bioengineering11101016>
5. Airlangga, G. (2024). Comparative Study of XGBoost, Random Forest, and Logistic Regression Models for Predicting Customer Interest in Vehicle Insurance. *Sinkron: Jurnal Dan Penelitian Teknik Informatika*, 8(4), 2542–2549. <https://doi.org/10.33395/sinkron.v8i4.14194>
6. Ghasemi, A., Hashtarkhani, S., Schwartz, D. L., & Shaban Nejad, A. Explainable artificial intelligence in breast cancer detection and risk prediction: A systematic scoping review. *Cancer Innovation*, 3:e136. <https://doi.org/10.1002/cai2.136>
7. Alelyani, T., Alshammari, M. M., Almuhan, A., & Asan, O. Explainable Artificial Intelligence in Quantifying Breast Cancer Factors: Saudi Arabia Context. *Healthcare*, 12(10):1025. <https://doi.org/10.3390/healthcare12101025>
8. Kutal, D. H., & Koseoglu, B. N. (2025). Breast Cancer Data Analysis Using Supervised Machine Learning Algorithms. *Cureus*, 17(10):e95011. <https://doi.org/10.7759/cureus.95011>
9. Darwich, M., & et al. (2024). An evaluation of the effectiveness of machine learning algorithms for breast cancer prediction. *International Journal of Medical Informatics*, 175:101550. <https://doi.org/10.1016/j.imu.2024.101550>
10. Anastasi, G., & et al. (2024). Machine learning techniques in breast cancer preventative analytics and feature selection effects. *Multimedia Tools and Applications*, 83:18775–xxxx. <https://doi.org/10.1007/s11042-024-18775-y>
11. Vrdoljak, J., Boban, Z., Barić, D., Šegvić, D., Kumrić, M., Avirović, M., & et al. (2023). Applying explainable machine learning models for detection of breast cancer lymph node metastasis in NST patients. *Cancers*, 15(3):634. <https://doi.org/10.3390/cancers15030634>
12. Naganandini, G., & Hulipalled, V. R. (2025). Breast cancer diagnosis using supervised machine learning with feature selection and hyperparameter tuning. *European Journal of Technical and Scientific Research*, 15(4):25634–25640. <https://doi.org/10.48084/etasr.11563>
13. La Moglia, A., & et al. (2025). Breast cancer prediction using machine learning classifiers: Feature selection and model comparison. *Bioinformatics Reports (Elsevier)*. <https://doi.org/10.1016/j.imu.2024.101550>
14. Lee, T. F., & et al. (2025). A machine learning model for predicting breast cancer risk and recurrence combining feature selection and ensemble methods. *Cancer Management and Research*, 17: S514693. <https://doi.org/10.2147/CMAR.S514693>
15. Islam, T., Sheakh, M. A., Tahosin, M. S., Hena, M. H., Akash, S., Jordan, Y. A. B., & et al. (2024). Predictive modelling for breast cancer classification with explainable AI. *Scientific Reports (Nature)*. <https://doi.org/10.1038/s41598-024-57740-5>
16. Ghasemi, A., Hashtarkhani, S., Schwartz, D. L., & Shaban-Nejad, A. (2024). Explainable artificial intelligence in breast cancer detection and risk prediction: A systematic scoping review. *Cancer Innovation*, 3(5), e136. <https://doi.org/10.1002/cai2.136>
17. Sadeghi, Z., Alizadehsani, R., Cifci, M. A., Kausar, S., Rehman, R., Mahanta, P., Bora, K., Almasri, A., Alkhawaldeh, R. S., Hussain, S., Alatas, B., Shoeibi, A., Moosaei, H., Hladik, M., Nahavandi, S., & Pardalos, P. M. (2024). A review of explainable artificial intelligence in healthcare. *Computers & Electrical Engineering*, 118, 109370. <https://doi.org/10.1016/j.compeleceng.2024.109370>
18. Prentzas, N., Kakas, A., & Pattichis, C. S. (2023). Explainable AI applications in the medical domain: A systematic review (arXiv:2308.05411). *arXiv*. <https://doi.org/10.48550/arXiv.2308.05411>