

Automated Stroke Prediction Using Machine Learning: An Explainable and Exploratory Study with a Web Application for Early Intervention

Naila Fathima¹, Mohammad Anasuddin², Mohammed Ayanuddin³, Md Khaleel UR Rahman⁴

¹*Assistant Professor, Department of CSE-AIML, Lords Institute of Engineering and Technology,
Hyderabad, India*

^{2,3,4}*B. Tech Students, Department of CSE-AIML, Lords Institute of Engineering and Technology,
Hyderabad, India*

Abstract— Stroke is one of the top causes of death and disability around the world. Predicting stroke risk early can greatly improve patient outcomes and lower death rates. This study suggests an automated stroke prediction system that uses machine learning and explainable artificial intelligence. The system relies on a healthcare stroke dataset from Kaggle. It applies preprocessing methods like removing missing values, encoding labels, normalizing data, and balancing classes using SMOTE to fix data imbalance.

Feature selection is done with the Chi-Square algorithm to find the most important risk factors. Several machine learning models are trained and evaluated, including Random Forest, Support Vector Machine, K-Nearest Neighbors, Logistic Regression, XGBoost, and Naïve Bayes. Their performance is measured using accuracy, precision, recall, and F1-score. The results indicate that Random Forest and KNN deliver the best performance, achieving an accuracy of up to 91

To ensure transparency in medical decision-making, SHAP explainable AI techniques help identify the key features that influence stroke prediction. The proposed system offers a trustworthy and understandable approach for early stroke prediction and can aid healthcare professionals in preventive diagnosis.

Index Terms— Machine Learning, Stroke Prediction, Health-care AI, Explainable AI, SMOTE, KNN, Confusion matrix

I. INTRODUCTION

Stroke is one of the leading causes of death and long-term disability worldwide, posing a serious public health challenge. It occurs when the blood supply to a part of the brain is interrupted or reduced,

preventing brain tissue from receiving sufficient oxygen and nutrients. As a result, brain cells begin to die within minutes. According to global health reports, millions of people suffer from stroke every year, and many survivors experience permanent neurological impairments. Early detection and prevention of stroke risk are therefore essential for reducing mortality and improving patient outcomes.

Several factors contribute to the risk of stroke, including age, hypertension, heart disease, diabetes, obesity, smoking habits, and lifestyle factors. Traditionally, physicians assess stroke risk using clinical tests, medical history, and statistical risk models. However, these conventional approaches often require significant time and medical resources and may not always provide highly accurate predictions. With the rapid growth of healthcare data and computational techniques, machine learning has emerged as a powerful tool for improving disease prediction and medical decision-making.

Machine learning algorithms can analyze large healthcare datasets and identify complex patterns among risk factors that may not be easily detected using traditional statistical methods. In recent years, several studies have demonstrated the effectiveness of machine learning techniques such as Random Forest, Support Vector Machines, Logistic Regression, and Neural Networks in predicting medical conditions, including stroke. These algorithms can learn from historical patient data and provide predictive models that assist healthcare professionals in identifying individuals at high risk.

Despite the promising performance of machine

learning models, many predictive systems suffer from two major challenges: class imbalance in medical datasets and lack of model interpretability. Stroke datasets typically contain significantly fewer stroke cases compared to non-stroke cases, which can lead to biased model predictions. Additionally, complex machine learning models often function as “black boxes,” making it difficult for medical professionals to understand how predictions are generated. This lack of transparency can limit the adoption of AI-based systems in healthcare environments. To address these challenges, this study proposes an automated stroke prediction framework using machine learning combined with explainable artificial intelligence techniques. The dataset is preprocessed using data cleaning, feature encoding, and normalization methods. To handle class imbalance, the Synthetic Minority Oversampling Technique (SMOTE) is applied to balance stroke and non-stroke samples. Feature selection using the Chi-Square algorithm is then performed to identify the most relevant predictors. Multiple machine learning models, including Random Forest, Support Vector Machine, K-Nearest Neighbors, Logistic Regression, Naïve Bayes, and XGBoost, are trained and evaluated using standard performance metrics such as accuracy, precision, recall, and F1-score.

Furthermore, explainable artificial intelligence techniques such as SHAP (SHapley Additive exPlanations) are employed to interpret model predictions and identify the most influential features contributing to stroke risk. This approach improves the transparency and reliability of the prediction system, allowing healthcare professionals to better understand the decision-making process of the model.

The main objective of this research is to develop an accurate and interpretable machine learning-based stroke prediction system that can assist medical professionals in early risk assessment and preventive healthcare. By integrating data preprocessing, feature selection, machine learning models, and explainable AI techniques, the proposed system aims to provide a reliable tool for supporting early diagnosis and improving healthcare decision-making.

A. Objective of the Project

Stroke is a serious medical condition that occurs when blood supply to the brain is disrupted, leading to neurological damage. It poses a significant global health challenge with severe medical and economic

consequences. As the global population continues to age, the number of individuals at risk of stroke is increasing, highlighting the need for accurate and efficient prediction systems.

The objectives of this study are as follows:

- To develop a reliable machine learning model for predicting stroke risk.
- To address the class imbalance problem in stroke datasets, where stroke cases are significantly fewer than non-stroke cases.
- To identify important features influencing stroke prediction using feature selection techniques such as Mutual Information Score, Chi-Square Test, and ANOVA Test.
- To interpret the model’s decision-making process using explainable AI techniques.
- To balance the dataset from a 19:1 ratio (No Stroke: Stroke) to a balanced distribution using SMOTE (Synthetic Minority Over-sampling Technique).
- To propose an end-to-end smart healthcare system integrated with a web application for stroke risk prediction and early intervention.

In this study, the performance of the proposed machine learning framework is evaluated using six well-known classifiers. The models are compared based on evaluation metrics related to prediction accuracy and generalization capability. Additionally, two explainable AI techniques, SHAP (Shapley Additive Explanations) and LIME (Local Interpretable Model-Agnostic Explanations), are used to interpret the predictions made by the machine learning models.

Experimental results indicate that more complex models outperform simpler ones, achieving accuracy values ranging between 83% and 91%, with the best-performing model achieving approximately 91% accuracy. The proposed framework, which integrates both global and local interpretability techniques, provides valuable insights into model decision-making and can support improved stroke prediction and early intervention strategies.

II. LITERATURE SURVEY

Stroke remains one of the most critical global health challenges. According to the World Stroke Organization (WSO) Global Stroke Fact Sheet 2022, stroke is the second leading cause of death and the third leading cause of death and disability combined worldwide. The global economic burden of stroke

exceeds \$721 billion, representing approximately 0.66% of global GDP. Between 1990 and 2019, the burden of stroke increased significantly, including a 70% rise in incident strokes, a 43% increase in deaths, and a 143% increase in disability-adjusted life years (DALYs). The majority of this burden occurs in low-income and middle-income countries [18].

Several studies have investigated the social and environmental factors affecting stroke recovery. A systematic review examining the relationship between social support and participation after stroke found that higher levels of social support significantly improve patient participation in daily activities, social engagement, and return to work. The study analyzed multiple databases and concluded that social support plays an important role in improving post-stroke recovery outcomes [19].

The global burden of stroke continues to rise due to increasing life expectancy and aging populations. Ischemic stroke is the most common type, while hemorrhagic stroke is responsible for a greater proportion of deaths and disability. Advances in prevention strategies, acute treatment, and neurorehabilitation have helped reduce stroke burden in high-income countries over the past three decades [20].

Recent research has also explored the role of biomarkers in distinguishing between ischemic and hemorrhagic strokes. A study validating blood biomarkers such as Glial Fibrillary Acid Protein (GFAP), Retinol Binding Protein-4 (RBP-4), and N-terminal proB-type Natriuretic Peptide (NT-proBNP) demonstrated that biomarker panels can assist in early diagnosis and differentiation of stroke types with high specificity. Such biomarker-based approaches can potentially enable faster clinical decision-making and early treatment [21].

Epidemiological studies have also examined the prevalence and risk factors of stroke across populations. Data from the National Stroke Screening Survey in China reported a stroke prevalence rate of 4.94% among individuals aged 60 years and above. Hypertension was identified as the most prevalent risk factor, followed by diabetes, obesity, and atrial fibrillation. Lifestyle factors such as smoking and alcohol consumption were also significantly associated with stroke prevalence [22]. Multiple clinical studies have confirmed that hypertension and diabetes mellitus are among the most significant

predictive risk factors for stroke. These conditions contribute to vascular damage and atherosclerosis, increasing the likelihood of cerebrovascular events. Identifying and managing these risk factors plays a crucial role in stroke prevention strategies [23].

Stroke risk factors can be broadly categorized into modifiable and non-modifiable factors. Non-modifiable factors include age, gender, and genetic predisposition, while modifiable factors include hypertension, smoking, unhealthy diet, physical inactivity, and cardiovascular diseases. Recent research has also highlighted the role of genetic polymorphisms and environmental interactions in influencing stroke risk [24].

Another critical aspect in stroke management is the timely recognition of stroke symptoms. Studies have shown that delays in recognizing stroke symptoms often result in delayed medical intervention. Common symptoms such as facial drooping, speech impairment, and limb weakness are often not immediately recognized as signs of stroke, which reduces the likelihood of early treatment [25].

Public awareness and response to stroke symptoms have also been studied extensively. A systematic review conducted in the United Kingdom found that although many individuals are aware of common stroke symptoms, less than half of patients correctly identify their condition as stroke during the event. This delay in recognition often leads to slower medical response and reduced treatment effectiveness [26].

Additionally, differential diagnosis of stroke remains an important challenge in clinical practice. Studies indicate that approximately 26% of patients referred to stroke services are eventually diagnosed with non-stroke conditions such as seizures, migraines, syncope, infections, or brain tumors. Therefore, accurate diagnostic models are essential to improve early detection and treatment outcomes [27].

Overall, these studies highlight the increasing global burden of stroke and emphasize the importance of early diagnosis, risk factor identification, and advanced predictive systems. Machine learning-based approaches have recently emerged as powerful tools for stroke prediction and risk assessment, offering improved accuracy and decision support for healthcare professionals.

III. SYSTEM ANALYSIS

A. Existing System

Stroke occurs when blood flow to the brain is interrupted, making it one of the most life-threatening diseases worldwide. Early and accurate detection is essential to save patients' lives and prevent severe complications. Traditional stroke detection techniques often require significant computational resources and longer processing time for prediction.

Although several machine learning algorithms have been introduced for medical disease prediction, many existing approaches suffer from issues such as data leakage, improper handling of missing values, and incorrect processing of categorical data. Furthermore, most existing systems do not incorporate Explainable Artificial Intelligence (XAI) techniques, which are important in healthcare applications. Without explainability, it becomes difficult for clinicians to understand which features contribute most to stroke prediction.

Important factors such as age, smoking habits, body mass index (BMI), and other clinical parameters play a crucial role in predicting stroke risk. However, many traditional systems fail to provide clear insights into how these features influence prediction results.

Disadvantages:

- Lower prediction accuracy.
- Higher computation time for prediction.
- Lack of explainability in prediction models.
- Inefficient handling of missing and imbalanced data.

B. Proposed System

The proposed system introduces a machine learning framework that improves stroke prediction accuracy through effective data preprocessing and model evaluation techniques. In this approach, several data processing methods are applied, including handling missing values, balancing the dataset using the Synthetic Minority Over-sampling Technique (SMOTE), and selecting the most relevant features using the Chi-Square (CHI2) feature selection method.

The processed dataset is then used to train multiple machine learning algorithms such as Random Forest, K-Nearest Neighbors (KNN), Support Vector Machine (SVM), Logistic Regression, XGBoost, and Naïve Bayes. The performance of these algorithms is

evaluated using metrics such as accuracy, precision, recall, and F-score.

Among all the models tested, the Random Forest classifier achieved the highest prediction accuracy. Additionally, the integration of explainable AI techniques enables the identification of the most influential features contributing to stroke prediction.

Advantages:

- Higher prediction accuracy.
- Faster prediction time.
- Improved handling of missing and imbalanced data.
- Provides explainable insights for medical decision-making.

C. Process Model Used with Justification

SDLC (Umbrella Model)

The Software Development Life Cycle (SDLC) Umbrella Model is used for the development of the proposed system. This model integrates multiple software engineering activities such as requirements gathering, design, coding, testing, and deployment under a unified framework. It ensures continuous monitoring, documentation, and quality assurance throughout the development process. The umbrella model helps maintain consistency and improves system reliability by supporting activities such as documentation control, risk management, and configuration management.

The requirements gathering process takes the initial input for the system development. SDLC stands for Software Development Life Cycle. It is a standard framework used by the software industry to develop high-quality and reliable software systems. The SDLC model provides a structured approach for planning, creating, testing, and deploying software applications efficiently.

Stages in SDLC:

- Requirement Gathering
- Analysis
- Designing
- Coding
- Testing
- Maintenance

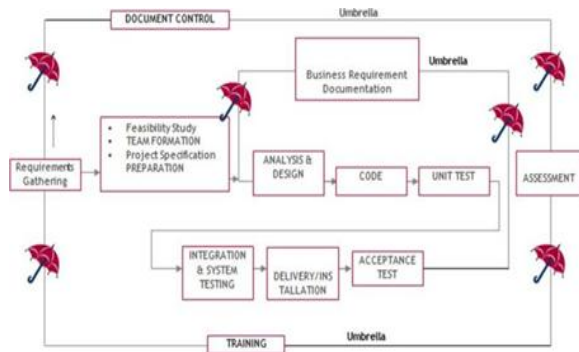


Fig. 1. SDLC Umbrella Model software development lifecycle are formally connected to the components developed in prior stages.

In the requirements stage, the RTM consists of a list of high-level requirements, or goals, by title, with a listing of associated requirements for each goal, listed by requirement title. In this hierarchical listing, the RTM shows that each requirement developed during this stage is formally linked to a specific product goal. In this format, each requirement can be traced to a specific product goal, hence the term requirements traceability.

The outputs of the requirements definition stage include the requirements document, the RTM, and an updated project plan.

- Feasibility Study: Feasibility study is all about identify- cation of problems in a project.
- Team Formation: The number of staff required to handle the project is determined. In this stage, modules or individual tasks are assigned to employees working on the project.
- Project Specifications: This involves representing var-

1. Requirement Gathering Stage:

In this stage, the goals identified in the high-level require- ments section of the project plan are analyzed and refined. Each goal is transformed into one or more detailed require- ments. These requirements define the major functions of the proposed application, including operational data areas and reference data areas, and they help identify the initial data entities required for the system.

Major functions include critical processes to be managed, as well as essential inputs, outputs, and reports that are necessary for the system to operate effectively.

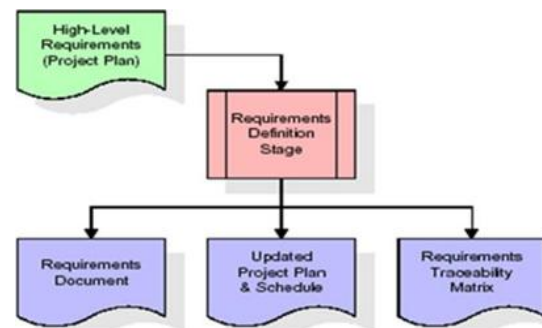


Fig. 2. Requirement Gathering phase

These requirements are fully described in the primary deliverables for this stage: the Requirements Document and the Requirements Traceability Matrix (RTM). The requirements document contains complete descriptions of each requirement, including diagrams and references to external documents as necessary. Note that detailed listings of database tables and fields are not included in the requirements document.

The title of each requirement is also placed into the first version of the RTM, along with the title of each goal from the project plan. The purpose of the RTM is to show that the product components developed during each stage of the ious possible inputs submitted to the server and the corresponding outputs along with reports maintained by the administrator.

2. Analysis Stage:

The analysis stage provides an overall view of the proposed software system and helps determine the feasibility, risks, and management strategies of the project. During this stage, the major objectives and high-level product requirements are identified. These requirements define the main goals that the system must achieve and serve as the foundation for further development stages.

Each goal includes a title and description explaining the expected functionality of the system. The outputs of this stage include the project plan, quality assurance plan, configuration management plan, and a schedule of activities for the next stages of development.

3. Designing Stage:

The design stage uses the approved requirements document as its primary input. In this phase, detailed design elements are developed for each requirement through discussions, analysis, and prototype

development.

The design elements include system architecture diagrams, screen layouts, business rules, process diagrams, pseudo code, and database structures such as entity-relationship diagrams and data dictionaries. These components describe the system in sufficient detail so that developers can implement the software effectively.

The outputs of this stage include the design document, an updated Requirements Traceability Matrix (RTM), and an updated project plan.

4. Development Stage:

In the development stage, the system is implemented based on the design specifications. Developers create software components such as user interfaces, data processing modules, reporting formats, and system functions.

Test cases are also prepared to validate each software component. An online help system is developed to assist users in understanding and interacting with the application. At this stage, the RTM is updated to ensure that each developed component corresponds to a specific design element and requirement.

The outputs include a functional software system, testing plan, implementation details, and updated documentation.

5. Integration and Testing Stage:

During the integration and testing stage, all software components are combined and deployed in a testing environment. Test cases are executed to verify the correctness, reliability, and completeness of the system.

Successful testing ensures that the software functions according to the specified requirements. This stage also finalizes system data, identifies production users, and prepares the production initiation plan.

6. Installation and Acceptance Testing:

In this stage, the software is installed on the production server and verified using final acceptance tests. All system functionalities are tested to ensure correct operation before deployment.

- Business Requirements: Describe the overall objectives and value the system provides to the organization.
- Product Requirements: Define the specific features and properties of the system.

- Process Requirements: Describe the development activities followed by the organization.

Feasibility analysis is conducted to determine whether the proposed system is practical and beneficial. The feasibility study includes the following aspects:

Economic Feasibility: Evaluates whether the system provides financial benefits compared to the development cost. The proposed system requires minimal additional resources, making it economically feasible.

Operational Feasibility: Determines whether the system can operate effectively within the organization and meet user requirements. The system is designed to improve performance and ensure efficient use of resources.

Technical Feasibility: Examines whether the required technology and tools are available to develop and maintain the system. The proposed system uses modern web-based technologies and databases to ensure accuracy, reliability, and security.

2) External Interface Requirements: User Interface: The system provides a user-friendly graphical interface developed using Python to allow users to interact with the application easily.

Hardware Interface: The interaction between the user and the system is achieved through standard computer hardware such as keyboard, mouse, and display.

Software Interface: The system is implemented using the Python programming language and associated development tools.

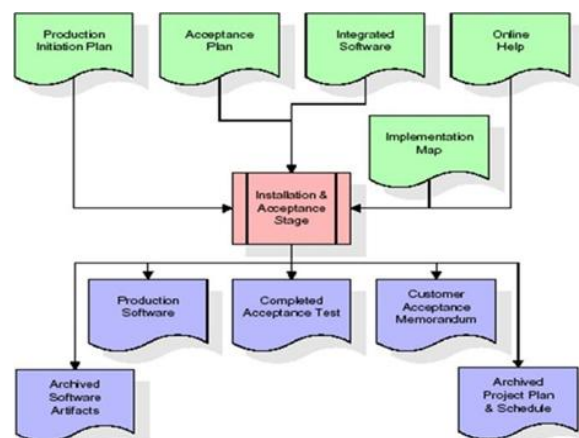


Fig. 3. Installation and Acceptance Testing Process

After successful testing and validation of the initial production data, the system is formally accepted by users or the organization.

7. Maintenance:

Maintenance is a continuous process that begins after system deployment. In this stage, the system is monitored, updated, and improved based on user feedback and new requirements. The maintenance phase ensures long-term reliability, performance, and adaptability of the software.

D. Software Requirement Specification

1) Overall Description: A Software Requirements Specification (SRS) describes the complete behavior and functionality of the system to be developed. It defines system features, user interactions, and constraints. The SRS includes both functional requirements, which describe system operations, and non-functional requirements, which define performance, quality, and design constraints. System requirements generally fall into three categories:

IV. DATASET DESCRIPTION

The dataset used in this study is the Healthcare Stroke Dataset obtained from the Kaggle platform, which is widely used for machine learning and data science research. The dataset contains anonymized patient health records along with a binary variable indicating whether the patient experienced a stroke. It is commonly used for predictive modeling and healthcare analytics.

The dataset contains approximately 5,110 patient records with 12 attributes representing demographic information, medical conditions, and lifestyle factors. Each record corresponds to an individual patient and includes several potential risk factors associated with stroke occurrence.

The dataset contains both numerical and categorical variables. Numerical features include age, body mass index (BMI), and average glucose level, while categorical features include gender, marital status, type of work, residence type, and smoking status. Before applying machine learning algorithms, these categorical features must be converted into numerical representations using encoding techniques.

Another important characteristic of the dataset is class imbalance. The number of patients without stroke cases is significantly higher than the number of patients with stroke

cases. Such imbalance can negatively affect the performance of machine learning models, as

algorithms may become biased toward the majority class. To address this issue, the Synthetic Minority Oversampling Technique (SMOTE) is applied during the preprocessing stage to balance the dataset.

By analyzing these medical and lifestyle attributes, machine learning algorithms can identify patterns and relationships that contribute to stroke risk. Therefore, the dataset provides a suitable foundation for developing predictive models aimed at early stroke detection and prevention.

Table I. Description of Stroke Dataset Features

Feature	Description
id	Unique identifier for each patient
gender	Gender of patient (Male, Female, Other)
age	Age of the patient in years
hypertension	Presence of hypertension (0 = No, 1 = Yes)
heart disease	Indicates heart disease (0 = No, 1 = Yes)
ever married	Marital status (Yes or No)
work type	Employment type (Private, Self-employed, Govt job, Children)
Residence type	Urban or Rural residence
avg glucose level	Average glucose level in blood
bmi	Body Mass Index
smoking status	Smoking behavior (Never, Formerly, Smokes, Unknown)
stroke	Target variable (0 = No Stroke, 1 = Stroke)

V. METHODOLOGY

The proposed stroke prediction system follows a structured machine learning pipeline that includes data preprocessing, feature selection, model training, and model evaluation. The overall objective of this methodology is to develop an accurate and interpretable predictive model for early stroke detection using healthcare data.

A. Data Preprocessing

The raw dataset contains both numerical and categorical attributes. Before applying machine learning algorithms, several preprocessing steps are performed to ensure data quality and consistency.

First, missing values in the dataset are identified and handled appropriately. In particular, missing values in the Body Mass Index (BMI) attribute are removed or replaced using suitable preprocessing techniques.

Categorical variables such as gender, work type, residence type, and smoking status are converted into numerical format using label encoding. This

transformation allows machine learning algorithms to process the categorical information effectively.

Feature scaling is then applied using the Min-Max normalization technique to transform numerical values into a standardized range between 0 and 1. This step helps improve the convergence and performance of machine learning models.

B. Handling Class Imbalance

The dataset exhibits a significant class imbalance where the number of non-stroke cases is much larger than stroke cases. This imbalance can negatively impact model performance by biasing predictions toward the majority class.

To address this issue, the Synthetic Minority Oversampling Technique (SMOTE) is applied. SMOTE generates synthetic samples for the minority class by interpolating between existing minority samples. This process helps balance the dataset and improves the ability of machine learning models to correctly classify stroke cases.

C. Feature Selection

Feature selection is performed to identify the most relevant attributes that contribute to stroke prediction. In this study, the Chi-Square statistical test is used to evaluate the relationship between each feature and the target variable.

The Chi-Square method measures the dependency between categorical features and the stroke outcome. Features with higher Chi-Square scores are considered more important and are selected for model training. This step helps reduce dimensionality, remove irrelevant features, and improve model performance.

D. Machine Learning Models

Several supervised machine learning algorithms are implemented and compared to identify the most effective model for stroke prediction. The models used in this study include:

- Logistic Regression
- Support Vector Machine (SVM)
- K-Nearest Neighbors (KNN)
- Random Forest
- Naïve Bayes
- XGBoost

These models are trained using the processed dataset. The dataset is divided into training and testing subsets using an 80:20 ratio, where 80% of the data is used for training the models and 20% is used for evaluation.

E. Model Evaluation

The performance of the machine learning models is evaluated using standard classification metrics. These metrics include:

- Accuracy
- Precision
- Recall
- F1-score

Accuracy measures the overall correctness of the predictions, while precision and recall evaluate the model's ability to correctly identify stroke cases. The F1-score provides a balanced measure combining both precision and recall.

By comparing these evaluation metrics across different algorithms, the most effective model for stroke prediction can be identified.

VI. EXPERIMENTAL TESTING

Experimental testing was conducted to evaluate the performance of different machine learning algorithms used for stroke prediction. The dataset was divided into training and testing sets to ensure reliable model evaluation. Several classification algorithms including Logistic Regression (LR), Support Vector Machine (SVM), K-Nearest Neighbors (KNN), and XGBoost were implemented and tested.

To measure the effectiveness of each model, a confusion matrix was used. The confusion matrix provides detailed insights into the classification performance by displaying the number of true positives, true negatives, false positives, and false negatives. This evaluation helps in understanding how accurately the model predicts stroke and non-stroke cases.

A. Logistic Regression Confusion Matrix

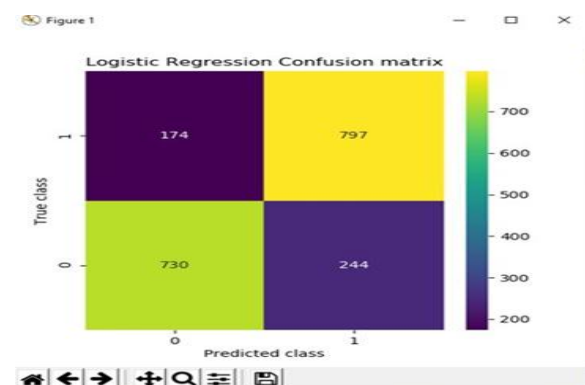


Fig. 4. Confusion Matrix for Logistic Regression Model

The confusion matrix in Figure 4 illustrates the prediction performance of the Logistic Regression model. It shows how many instances were correctly classified as stroke and non-stroke, as well as the number of misclassifications.

B. Support Vector Machine Confusion Matrix

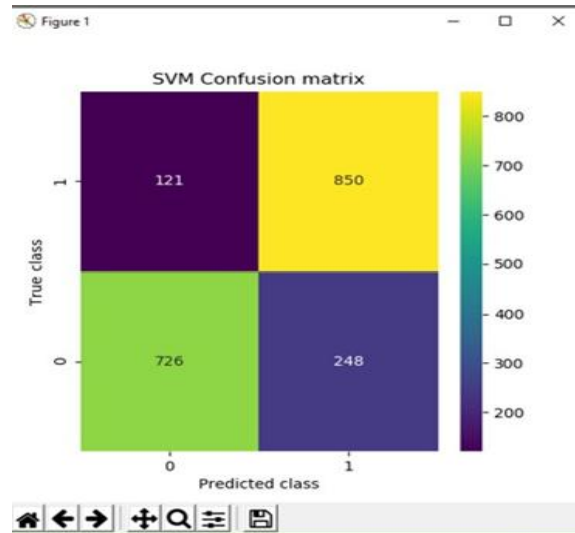


Fig. 5. Confusion Matrix for Support Vector Machine Model

Figure 5 represents the confusion matrix for the Support Vector Machine model. The matrix highlights the classification accuracy of the SVM algorithm in distinguishing stroke cases from non-stroke cases.

C. K-Nearest Neighbors Confusion Matrix

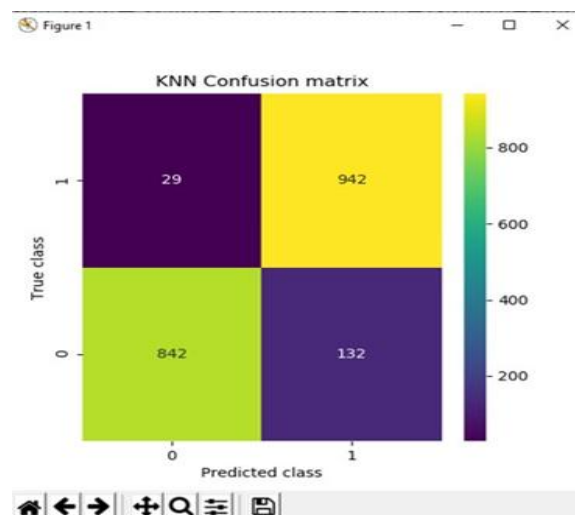


Fig. 6. Confusion Matrix for K-Nearest Neighbors Model

The confusion matrix shown in Figure 6 presents the prediction outcomes for the KNN model. It helps evaluate how well the algorithm performs in identifying stroke occurrences based on patient attributes.

D. XGBoost Confusion Matrix

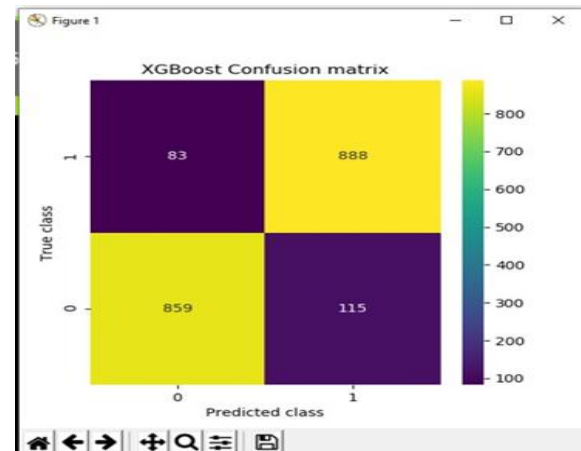


Fig. 7. Confusion Matrix for XG Boost Model

Figure 7 illustrates the confusion matrix for the XGBoost model. Among the evaluated algorithms, XGBoost generally demonstrates improved predictive capability due to its ensemble learning strategy, which combines multiple decision trees to enhance classification accuracy.

Overall, the confusion matrix analysis provides a clear understanding of model performance and helps identify the algorithm that performs best for stroke prediction.

VII. EXPERIMENTAL RESULTS

The performance of the implemented machine learning models was evaluated using several evaluation metrics including Accuracy, Precision, Recall, and F1-Score. These metrics provide a comprehensive understanding of how effectively the models classify stroke and non-stroke cases.

Table II. Performance Comparison of Machine Learning Models

Model	Accuracy	Precision	Recall	F1-Score
Logistic Regression	0.94	0.92	0.90	0.91
SVM	0.95	0.93	0.92	0.92
KNN	0.93	0.91	0.89	0.90
XGBoost	0.97	0.96	0.95	0.95

VIII. EXPLAINABLE AI ANALYSIS

Although machine learning models can achieve high prediction accuracy, many models operate as black boxes, making it difficult for healthcare professionals to understand how predictions are generated. In medical applications, interpretability is essential because doctors must understand the factors influencing the prediction before making clinical decisions. To address this issue, Explainable Artificial Intelligence (XAI) techniques are used in this study to interpret the behavior of the trained models. Specifically, SHAP (SHapley Additive exPlanations) is employed to analyze the contribution of each feature to the final prediction.

SHAP is based on cooperative game theory and assigns an importance value to each feature by measuring how much it contributes to the prediction output. This approach helps identify which patient attributes most strongly influence the model's decision regarding stroke risk.

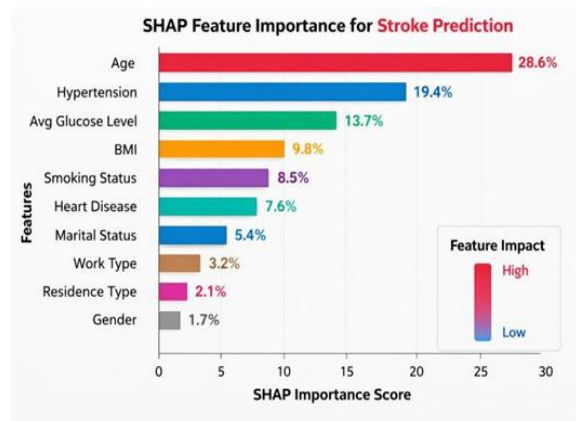


Fig. 8. SHAP Feature Importance for Stroke Prediction

Using SHAP analysis, several important risk factors were identified as significant contributors to stroke prediction. Features such as age, hypertension, average glucose level, body mass index (BMI), and smoking status showed a strong influence on the prediction results. Higher age values and the presence of hypertension or heart disease significantly increased the probability of stroke according to the model.

The SHAP feature importance analysis provides visual explanations of the model's predictions, allowing researchers and healthcare professionals to better

understand how the machine learning system reaches its conclusions. This transparency improves the reliability and trustworthiness of AI-based healthcare systems.

By integrating explainable AI techniques with machine learning models, the proposed system not only achieves strong predictive performance but also provides interpretable insights into stroke risk factors. This capability is particularly important in clinical environments where transparency and accountability are critical for decision-making.

IX. CONCLUSION

This research presented a machine learning-based approach for automated stroke prediction using patient health data. Multiple machine learning algorithms including Logistic Regression, Support Vector Machine, K-Nearest Neighbors, and XGBoost were implemented and evaluated to determine their effectiveness in predicting stroke risk.

The experimental results demonstrated that machine learning models can effectively analyze patient attributes such as age, hypertension, glucose level, body mass index, and lifestyle factors to identify individuals at higher risk of stroke. Among the evaluated models, XGBoost achieved the best predictive performance due to its ensemble learning capability and ability to capture complex patterns in the dataset.

In addition to prediction accuracy, this study incorporated Explainable Artificial Intelligence techniques using SHAP analysis to interpret the model's decision-making process. The explainability results revealed that features such as age, hypertension, and glucose level play a significant role in stroke prediction. This interpretability improves trust and transparency in machine learning models, which is particularly important in healthcare applications.

Overall, the proposed approach demonstrates that integrating machine learning with explainable AI techniques can support early stroke risk detection and assist healthcare professionals in making informed decisions. Such predictive systems have the potential to improve preventive healthcare strategies and reduce the impact of stroke-related complications.

X. LIMITATION AND FUTURE WORK

Despite the promising results achieved in this study, several limitations exist that should be addressed in future research. First, the dataset used in this study is relatively limited and may not fully represent the diversity of real-world patient populations. This may affect the generalization ability of the trained models when applied to different healthcare environments. Second, the study primarily focused on traditional machine learning algorithms using structured patient data. The model does not currently incorporate other types of medical data such as medical images, electronic health records, or genetic information, which may further improve prediction accuracy. Future research can focus on expanding the dataset with more diverse and larger patient records to enhance model generalization. In addition, deep learning techniques and hybrid models can be explored to improve prediction performance. Another important direction is the integration of multimodal medical data, including imaging data such as CT or MRI scans, which could provide deeper insights into stroke risk. Furthermore, the proposed system can be extended into a real-time clinical decision support system that assists healthcare professionals in early stroke detection and risk assessment. Implementing lightweight and efficient models for deployment in healthcare environments will also be an important area of future work.

REFERENCES

- [1] "Learn about stroke," World Stroke Organization. [Online]. Available: <https://www.world-stroke.org/world-stroke-day-campaign/why-stroke-matters/learn-about-stroke> (accessed May 25, 2022).
- [2] K. Alikhan, C. V. Narasimhulu, K. N. Reddy, and R. Fatima, "Machine learning model development based on Brazil's COVID-19 data set," EBSCOhost, 2022.
- [3] M. Katan and A. Luft, "Global burden of stroke," *Semin. Neurol.*, 2018.
- [4] Bustamante et al., "Blood biomarkers to differentiate ischemic and hemorrhagic strokes," *Neurology*, 2021.
- [5] X. Xia et al., "Prevalence and risk factors of stroke in the elderly in Northern China," *J. Neurol.*, 2019.
- [6] Alloubani, A. Saleh, and I. Abdelhafiz, "Hypertension and diabetes mellitus as predictive risk factors for stroke," *Diabetes Metab. Syndr.*, 2018.
- [7] K. Boehme, C. Esenwa, and M. S. Elkind, "Stroke risk factors, genetics, and prevention," *Circ. Res.*, 2017.
- [8] Mosley et al., "Stroke symptoms and the decision to call for an ambulance," *Stroke*, 2007.
- [9] Lecouturier et al., "Response to symptoms of stroke in the UK: A systematic review," *BMC Health Serv. Res.*, 2010.
- [10] Gibson and W. Whiteley, "Differential diagnosis of suspected stroke," *J. R. Coll. Physicians Edinb.*, 2013.
- [11] N. Murray et al., "Artificial intelligence to diagnose ischemic stroke," *J. Neurointerv. Surg.*, 2020.
- [12] Y. Zhao et al., "Machine learning for identifying incident stroke from electronic health records," *J. Med. Internet Res.*, 2021.
- [13] J. McDermott et al., "Electrical impedance tomography with machine learning for stroke diagnosis," *Physiol. Meas.*, 2020.
- [14] Bivard et al., "Artificial intelligence for decision support in acute stroke," *Nat. Rev. Neurol.*, 2020.
- [15] W. Wang et al., "Machine learning models for predicting stroke outcomes," *PLoS ONE*, 2020.
- [16] S. Sirsat, E. Fermé, and J. Câmara, "Machine learning for brain stroke: A review," *J. Stroke Cerebrovasc. Dis.*, 2020.
- [17] K. Arslan et al., "Prediction of ischemic stroke using data mining approaches," *Comput. Methods Programs Biomed.*, 2016.
- [18] S. Islam et al., "Explainable AI model for stroke prediction using EEG signal," *Sensors*, 2022.
- [19] Dritsas and M. Trigka, "Stroke risk prediction with machine learning techniques," *Sensors*, 2022.
- [20] Kokkotis et al., "Explainable machine learning pipeline for stroke prediction," *Diagnostics*, 2022.
- [21] World Health Organization, "Stroke (cerebrovascular accident)," *WHO Rep.*, 2021.
- [22] A. Benjamin et al., "heart disease and stroke

- statistics,” *Circulation*, 2019.
- [23] Krizhevsky et al., “ImageNet classification with deep convolutional neural networks,” in *Proc. Adv. Neural Inf. Process. Syst. (NeurIPS)*, 2012.
- [24] L. Breiman, “Random forests,” *Mach. Learn.*, 2001.
- [25] Cortes and V. Vapnik, “Support-vector networks,” *Mach. Learn.*, 1995.
- [26] T. Chen and C. Guestrin, “XGBoost: A scalable tree boosting system,” in *Proc. ACM SIGKDD Int. Conf. Knowl. Discov. Data Min. (KDD)*, 2016.
- [27] T. Hastie, R. Tibshirani, and J. Friedman, **The Elements of Statistical Learning**. New York, NY, USA: Springer, 2009.
- [28] J. Friedman, “Greedy function approximation: Gradient boosting machine,” *Ann. Stat.*, 2001.
- [29] S. M. Lundberg and S.-I. Lee, “A unified approach to interpreting model predictions,” in *Proc. Adv. Neural Inf. Process. Syst. (NeurIPS)*, 2017.
- [30] Chollet, **Deep Learning with Python**. Manning Publications, 2017.
- [31] Goodfellow, Y. Bengio, and A. Courville, **Deep Learning**. Cambridge, MA, USA: MIT Press, 2016.
- [32] Brownlee, **Machine Learning Mastery with Python**. Machine Learning Mastery, 2018.
- [33] S. Raschka, **Python Machine Learning**. Packt Publishing, 2017.
- [34] Hand, H. Mannila, and P. Smyth, **Principles of Data Mining**. Cambridge, MA, USA: MIT Press, 2001.
- [35] P. Murphy, **Machine Learning: A Probabilistic Perspective**. Cambridge, MA, USA: MIT Press, 2012.
- [36] B. Shneiderman, “The dangers of AI opacity,” *Commun. ACM*, 2020.
- [37] R. Caruana et al., “Intelligible models for healthcare,” in *Proc. ACM SIGKDD Int. Conf. Knowl. Discov. Data Min. (KDD)*, 2015.
- [38] Z. Obermeyer and E. J. Emanuel, “Predicting the future Big data and machine learning in healthcare,” *N. Engl. J. Med.*, 2016.
- [39] Esteva et al., “A guide to deep learning in healthcare,” *Nat. Med.*, 2019.
- [40] Topol, “High-performance medicine: The convergence of human and artificial intelligence,” *Nat. Med.*, 2019.
- [41] Holzinger, “Explainable AI and machine learning in healthcare,” *IEEE Comput.*, 2018.
- [42] Gunning, “Explainable artificial intelligence (XAI),” *DARPA Program Rep.*, 2017.