

Predicting Medication-Related Risk in Chronic Disease Management Using Machine Learning

Blessing Malcolm Awukam¹, Chinaza Felicia Nwakobe², Sunday Okafor³

¹School of Business, Technology, and Health Care Administration, Capella University, Minnesota, United States.

²Department of Emergency Medicine, Betsi Cadwaladr University Health Board, Bangor, United Kingdom.

³Media Department, Hochschule Darmstadt University of Applied Sciences, Darmstadt, Germany.

Received: 01 October 2025 Revised: 05 October 2025 Accepted: 09 October 2025 Published: 14 October 2025

Abstract - The application of Machine Learning (ML) in chronic disease management has emerged as a transformative approach to improving medication safety and therapeutic outcomes. By leveraging predictive models, healthcare systems can proactively identify patients at heightened risk of medication-related complications, such as adverse drug events, polypharmacy interactions, and inappropriate dosing. This study explores the role of machine learning-driven predictive modelling in mitigating medication-related risks within chronic disease care, emphasising its potential to enhance clinical decision-making, reduce hospital readmissions, and improve patient safety. Key challenges, including data heterogeneity, model interpretability, and integration with electronic health records, are critically examined, alongside opportunities for refining prediction accuracy and optimising personalised medication management. The paper underscores the necessity of interdisciplinary collaboration among clinicians, data scientists, and health informaticians to fully harness the potential of machine learning in medication risk prediction. Such collaboration is vital to advancing precision medicine and ensuring safer, more effective therapeutic strategies for patients with chronic diseases.

Keywords - Machine Learning, Predictive Modelling, Medication Safety, Chronic Disease Management, Adverse Drug Events, Polypharmacy, Healthcare Informatics.

I. INTRODUCTION

Chronic diseases such as diabetes, hypertension, cardiovascular disorders, and chronic kidney disease represent a growing global health burden, accounting for over 70% of worldwide mortality [1]. Pharmacotherapy is a key to the successful long-term management of such conditions. Nevertheless, multidrug therapy and its intricacy particularly in multimorbid patients are likely to amplify the risks associated with medication including adverse drug events (ADEs), drug-drug interactions (DDIs), dosing mistakes, and non-adherence [2]. These are significant risks in terms of hospitalisations, morbidity and healthcare expenditure, and thus, proactive medication risk assessments strategies are in dire need [3].

Conventional methods of risk management in medication, through manual review of charts, or rule-based identifier systems are usually not adequate because of their use of fixed levels and their inability to be adjusted to the unique patient characteristics [4]. On the contrary, machine learning suggests dynamic and data-driven approaches that are capable of yielding nonlinear relationships between demographic, clinical, and pharmacological variables. ML models can then predict the risk level of individual patients with ADEs or ineffective treatment using large-scale health data such as electronic health records (EHRs) prescription history, and laboratory outcomes [5].

As an example, ML algorithms have the potential to examine the interactions of medications, comorbidities and laboratory trends and identify patients with high risk of renal toxicity, bleeding or glycemic fluctuations [6]. Such predictive insights enable clinicians to make individualised changes to medication in the management of chronic diseases to enhance outcomes of the therapy and reduce harm. Furthermore, the latter models may be used to support clinical decision support systems (CDSS), which will deliver real-time notifications about possible

medication errors and streamline treatment paths. Although the potential of ML in predicting risk of medications is increasing, a number of issues remain. The quality of data and its interoperability are also major challenges with the healthcare datasets frequently having gaps in values, inconsistency, and partial records of patients [7]. As well, black box features of most ML models decrease interpretability of clinical models in clinics, which may reduce clinician trust and adoption [8]. Ethical and regulatory concerns especially those touching on patient privacy, algorithmic bias and accountability also require a keen consideration before full scale implementation [9]. To reach the potential of ML in minimising the risks related to medications, strong, interpretable, and clinically validated predictive frameworks are required. Clinicians, data engineers, and policymakers should work closely together in order to integrate ML models into the processes of managing chronic diseases. ML-based risk prediction of medication can help considerably develop precision pharmacotherapy, improve patient safety, and lower healthcare costs with proper infrastructure and governance [10].

II. KEY CONSIDERATIONS FOR APPLYING MACHINE LEARNING TO MEDICATION RISK PREDICTION

The successful use of machine learning for medication-risk prediction in chronic disease management depends on more than algorithmic performance. A model can only be considered valuable for clinical use when it is reliable, interpretable, adaptable, and easily integrated in healthcare workflows. In this section, the author defines five critical factors in the development and deployment of reliable predictive systems in medication safety management.

A. Model Generalisability and Validation

Model generalisability guarantees predictive accuracy to be consistent in varying clinical settings as well as patient groups. A model that has been trained in one institution might not work similarly in a different institution since the demographics can vary, the prescribing practice can vary, or the data-recording practices can change. Making extensive validation is hence essential. The model is validated externally by a variety of datasets and healthcare environments with the ability to retain its accuracy, precision, and recall in new settings. Calibration checks should also be done to confirm that the predicted probabilities are an indicator of real clinical outcomes [11]. Transparent documentation of validation methods and results strengthens scientific credibility and supports regulatory acceptance. A well-validated model promotes clinician trust and ensures that predictive systems remain dependable when applied in real-world care.

B. Clinical Integration and Workflow Alignment

For predictive models to influence clinical decision-making effectively, they must integrate seamlessly into existing healthcare workflows. Even the most accurate model is of limited value if its outputs are not accessible or usable at the point of care [12]. Integration with electronic health record systems enables real-time support, such as automated alerts for unsafe drug combinations, dosing errors, or patient-specific contraindications. However, excessive or poorly designed notifications can cause alert fatigue, reducing clinical attention to critical risks. Collaboration among clinicians, pharmacists, and technology specialists during system design helps ensure that predictive insights are clinically relevant, appropriately prioritised, and presented in an actionable format. Proper workflow alignment turns predictive analytics into a practical decision-support asset rather than a technical burden.

C. Sustainability and Continuous Learning

Machine learning models must evolve alongside changes in medical practice, drug availability, and population health trends. Continuous monitoring, retraining, and recalibration ensure that predictive accuracy and clinical relevance are preserved over time. As new patient data and medication records accumulate, incremental retraining allows models to adapt without complete redevelopment. Ongoing surveillance of performance metrics helps detect data drift or degradation early, prompting timely updates. Sustainability also requires structured documentation and version control to track each change in model architecture, dataset composition, and parameter configuration. Embedding continuous learning into model governance safeguards long-term reliability and regulatory transparency [13].

D. Data Integrity and Accessibility

Data quality remains the foundation of every predictive model. In medication-risk assessment, healthcare data are often sourced from multiple systems such as electronic health records, pharmacy databases, laboratory information systems, and patient-reported outcomes, each with its own standards and completeness levels. Inconsistent documentation, missing values, and a lack of standardised medication coding can seriously reduce model accuracy [14]. For example, omission of over-the-counter drugs or herbal supplements may distort risk estimation, while irregular laboratory entries can obscure emerging adverse trends such as renal impairment. Data fragmentation across healthcare providers further limits reliability. Patients with chronic diseases often receive treatment from multiple specialists, leading to scattered records that hinder long-term tracking. Implementing interoperable data infrastructures, enforcing strict governance standards, and promoting structured medication data formats are therefore essential. Incorporating real-world evidence, such as pharmacy refill histories and wearable device data, can enrich models and improve predictive precision [15].

E. Model Interpretability

Interpretability remains central to clinical adoption. Healthcare providers must understand how and why a model produces specific predictions before trusting its recommendations. Advanced algorithms, particularly ensemble and deep-learning techniques, often function as opaque systems that lack transparent explanations. If a model identifies a patient as high risk without clear reasoning, clinicians may hesitate to act. In safety-critical environments, interpretability is not optional; it is necessary for accountability, patient confidence, and informed decision-making [16,17]. To address this, interpretable modelling frameworks and post-hoc explanation methods are increasingly used to show which factors most strongly influence each prediction. This transparency bridges the gap between data science and clinical reasoning, supporting responsible and evidence-aligned use of predictive analytics in medication management [18].

III. INTEGRATION OF MACHINE LEARNING MODELS WITH EXISTING HEALTHCARE SYSTEMS

The successful implementation of ML models for medication risk prediction depends heavily on their effective integration into existing healthcare infrastructures. Despite the growing sophistication of predictive algorithms, their impact remains limited unless seamlessly embedded within clinical workflows, EHRs, and decision support systems. However, integration faces substantial challenges arising from infrastructure limitations, data silos, and workflow misalignment, particularly in healthcare settings with constrained resources [19]. Many hospitals and outpatient clinics continue to rely on legacy EHR systems that are not optimised for interoperability or real-time analytics. These systems often store data in fragmented formats, making it difficult to incorporate ML-driven medication risk predictions directly into point-of-care decision-making [20]. For instance, an ML model identifying patients at high risk of ADEs may not be able to automatically trigger alerts within an EHR if the underlying software lacks standardised interfaces or APIs. As a result, critical predictive insights remain underutilised, reducing the model's potential to prevent harm.

Integration of ML systems with the current clinical tools poses technical and operational problems even in technologically developed settings. The design of how to align model outputs with the workflow of physicians must be considered to prevent the occurrence of alert fatigue or the disruption of workflow. In case predictive alerts are not timely or too sensitive, clinicians can ignore them, which destroys the functionality of a system. The key to automation and clinical judgment is to strike an optimal balance that ensures acceptance and continued use [21]. In order to address these impediments, healthcare organisations should give interoperability of systems and access to cross-platforms data integration a top priority. The implementation of standardised data exchange models, e.g. HL7 FHIR (Fast Healthcare Interoperability Resources) can help to ensure the smooth exchange of data between the ML models, EHRs, and pharmacy systems. In addition, it can be integrated with CDSS to produce personalised medication alerts in real-time, including recognising unsafe drug interactions, dosing mistakes, and organ toxicity depending on patient-specific factors [22]. To make sure that ML models are contextually relevant and can be acted upon by clinicians, collaboration between data scientists and software vendors with healthcare providers is essential. As an illustration, the predictive models may be incorporated into the chronic care management platform, which would facially enable automatic identification of potentially at-

risk patients due to medication non-adherence or polypharmacy complications. The role of governments and regulatory authorities is also crucial since they can fund, develop digital health standards, and enforce regulations on data security that would help to deploy ML responsibly [23].

Public-private partnerships can further accelerate this integration by combining clinical expertise with technological innovation. These partnerships have the potential to facilitate scaled solutions based on the specific situation of the healthcare at the local level and make sure that predictive systems meet clinical validation criteria and ethical principles [24]. It all depends on the fact that implementing the concept of ML-based medication risk forecasting in the healthcare system is not only a technical task but also an organisational change that needs a certain investment in infrastructure, policy adjustment, and multidisciplinary co-work to achieve its potential in enhancing patient safety and chronic disease outcomes.

IV. IMPLEMENTATION OF PREDICTIVE MODELING

The implementation of predictive modelling of medication-related risk in chronic disease management is done in a structured and systematic pipeline that is aimed at ensuring accuracy, reproducibility, and clinical relevance. It has six steps; data collection, preprocessing, feature selection, data splitting, model development and model evaluation.

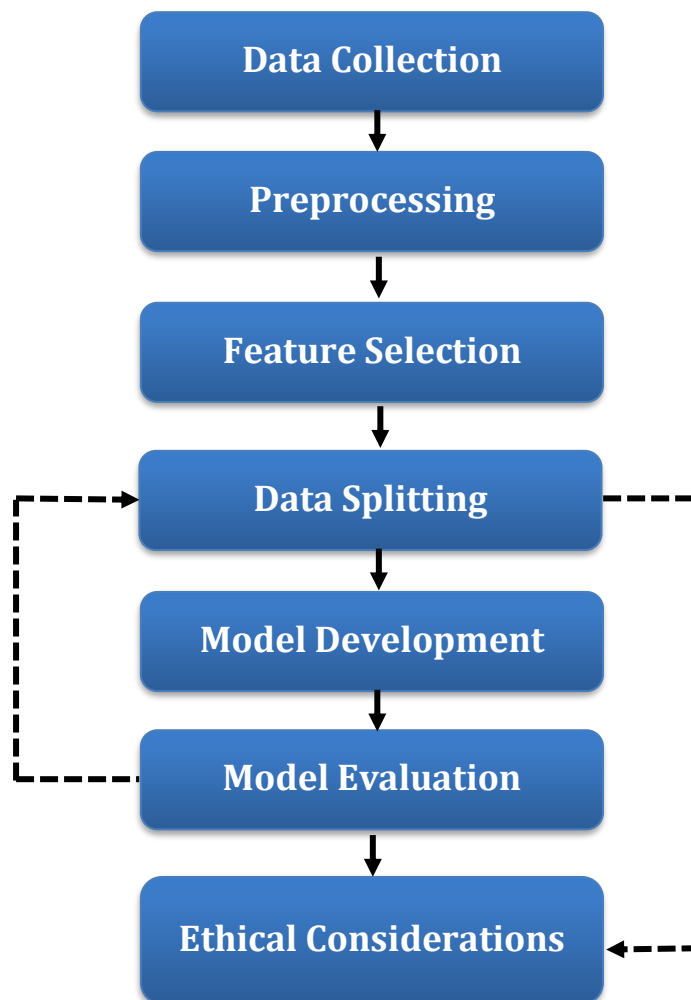


Figure 1. The Predictive Modelling Steps

This workflow, as shown in Figure 1, is the basis of creating dependable ML-based risk prediction systems, which can detect patients at increased risk of adverse drug events (ADEs), polypharmacy comorbidities, and non-adherence to medicine. All stages play an essential part in enhancing model performance and providing predictive insights that are consistent with clinical practice in the real world.

V. DATA COLLECTION

The dataset used for modelling medication-related risk comprises synthetic patient data representing individuals with chronic conditions such as diabetes, hypertension, and heart failure. Each patient record contains a combination of clinical, pharmacological, and demographic variables relevant to medication safety assessment.

These variables provide insight into individual treatment profiles and potential risk factors for adverse outcomes. The key features include:

- **Age:** Patient's chronological age (years).
- **Gender:** Categorical variable (Male/Female).
- **Comorbidities:** Categorical list of co-existing chronic conditions (e.g., Diabetes, Chronic Kidney Disease, Hypertension).
- **Medication Count:** Number of concurrent prescribed medications (integer).
- **High-Risk Drug Use:** Binary variable indicating exposure to drugs associated with elevated ADE risk (e.g., anticoagulants, NSAIDs, opioids).
- **Renal Function (eGFR):** Continuous variable representing estimated glomerular filtration rate, indicative of kidney function.
- **Liver Enzyme Levels (ALT/AST):** Continuous variable indicating hepatic function relevant to drug metabolism.
- **Drug-Drug Interaction (DDI) Flag:** Binary indicator showing presence or absence of potential DDI based on known pharmacological profiles.
- **Medication Adherence:** Ordinal variable (High, Medium, Low) based on pharmacy refill or digital adherence metrics.
- **Adverse Drug Event Risk:** The target variable, representing overall medication-related risk, is classified as Low, Moderate, or High.

Table 1: Sample of Dataset Used for Predictive Modelling in Medication Risk Assessment

Age	Gender	Comorbidities	Medication Count	High-Risk Drug Use	eGFR	ALT	DDI Flag	Adherence	Risk Level
65	Male	Diabetes, HTN	8	Yes	48	52	Yes	Medium	High
52	Female	None	3	No	92	31	No	High	Low
74	Male	CKD, CHF	10	Yes	36	44	Yes	Low	High
59	Female	HTN	5	No	72	37	No	Medium	Moderate
68	Male	Diabetes	7	Yes	55	46	Yes	High	High
47	Female	Asthma	4	No	88	40	No	High	Low

Note: The Dataset contains 120 synthetic records with 10 features and 1 target variable.

VI. DATA PREPROCESSING

Before training any machine-learning model, raw clinical and pharmacological data must be cleaned and standardised to ensure accurate, unbiased, and reproducible predictions. Data preprocessing transforms heterogeneous patient information into a structured format that algorithms can interpret effectively. The following steps, categorical encoding, handling of missing values, and scaling of numerical features, were applied to the medication-risk dataset to enhance data integrity and model performance.

A. Conversion of Categorical Variables

Several features in the dataset, such as Gender, Comorbidities, Medication Adherence, and High-Risk Drug Use, are categorical in nature. These were converted into numerical representations using Label Encoding and One-Hot Encoding techniques. This transformation allows machine-learning algorithms particularly tree-based and linear models to interpret categorical information quantitatively. The adherence (High, Medium, Low) were, for example, converted into ordinal (2, 1, 0) so that the natural order is maintained, and binary (High-Risk Drug Use: Yes/No) variables were converted into 1 and 0, respectively. The encoding of these variables is necessary to ensure the preservation of clinical meaning, and the lack of a consistent encoding of these variables may lead to

the distortion of relationships between drug exposure and adverse outcomes. With this step, the dataset can be converted to be machine-readable but still has its semantic structure.

B. Handling Missing Values

Incomplete or missing clinical data is a pervasive issue in healthcare analytics. Missing entries such as absent lab results (eGFR, ALT/AST) or incomplete medication lists can introduce bias and reduce model accuracy. To overcome this, a hybridic approach of imputation was used. The median or the mean as a substitution of continuous variables was done based on the distribution, and mode values were substituted on the categorical fields. Where some of the predictors of interest such as renal or hepatic functioning were not provided, the corresponding patient record was not included in playing the model to ensure integrity in the dataset. This method reduces distortion and leaves enough sample size to perform adequate analysis.

C. Scaling Numerical Features

Numerical attributes, including Age, Medication Count, eGFR, and ALT/AST, were normalised using the StandardScaler method, which rescales values to have zero mean and unit variance. Scaling is used to make sure that the advantages of all numerical variables are proportional in the optimisation of the model. The large variables (such as Medication Count) will overwhelm smaller variables (such as Liver Enzyme Levels) otherwise, and results in biased weight distribution, and unstable gradients in model-based systems like support-vector machines or neural networks. This preprocessing can balance the scales of all continuous predictors, thus ensuring one gets a balanced process of learning and improving the interpretability and stability of further modelling.

VII. FEATURE SELECTION

Feature selection is a crucial phase in developing reliable and interpretable medication-risk prediction models. It involves identifying the most influential predictors of ADEs and filtering out irrelevant or redundant variables that may introduce noise or bias. This process not only enhances model performance but also improves interpretability, enabling clinicians to understand the clinical rationale behind each prediction. In this paper, domain knowledge along with statistical correlation analysis and machine learning-based feature importance ranking were combined in order to identify the most applicable predictors of medication-related risk in the management of chronic diseases. Preliminary models did not eliminate any of the ten variables to determine baseline performance. This was followed by feature-selection methods of higher sophistication to narrow down the input space and only the most meaningful factors were retained.

A. Correlation and Redundancy Analysis

Pearson correlation heatmap was produced to determine the relationships between continuous variables (Age, Medication Count, eGFR and Liver Enzyme Levels). This was done to determine the possible multicollinearity, and thus to make sure that the very high-correlated variables did not corrupt the interpretation of the model and decrease the generalizability. Categorical variables were also assessed on the basis of the Cramer V statistic that is used to measure the relationship between nominal variables such as Comorbidities and High-Risk Drug Use. Figure 2 represents the correlation framework of the dataset with identification of the relations between features and their possible role in medication-related risk.

B. Feature Importance Ranking

The relative importance of each feature in predicting the Risk Level target variable was ranked with the help of machine learning algorithms, namely, Random Forest and Gradient Boosting Classifiers. The models under consideration have impurity reduction and information gain measures in their contributions. The highest ranked features were common to all methods which supports their usefulness in predicting medication risks.

The selected features include:

- **Medication Count:** A strong predictor of polypharmacy-related risk and dosing complexity.
- **High-Risk Drug Use:** Indicates potential exposure to drugs commonly associated with ADEs.
- **eGFR:** Reflects renal function, which directly affects drug metabolism and clearance.
- **Comorbidities:** Represent underlying disease states influencing drug response and sensitivity.

- **Medication Adherence:** Captures behavioural and compliance factors that modulate treatment outcomes.
- **Age:** A demographic determinant linked to altered pharmacokinetics and pharmacodynamics.

These predictors were selected based on their clinical relevance, statistical significance, and model-derived importance scores, ensuring that the dataset remains both scientifically grounded and computationally efficient. By narrowing the feature space to these key variables, the final model achieves improved interpretability and reduced overfitting while maintaining predictive power across heterogeneous patient populations.

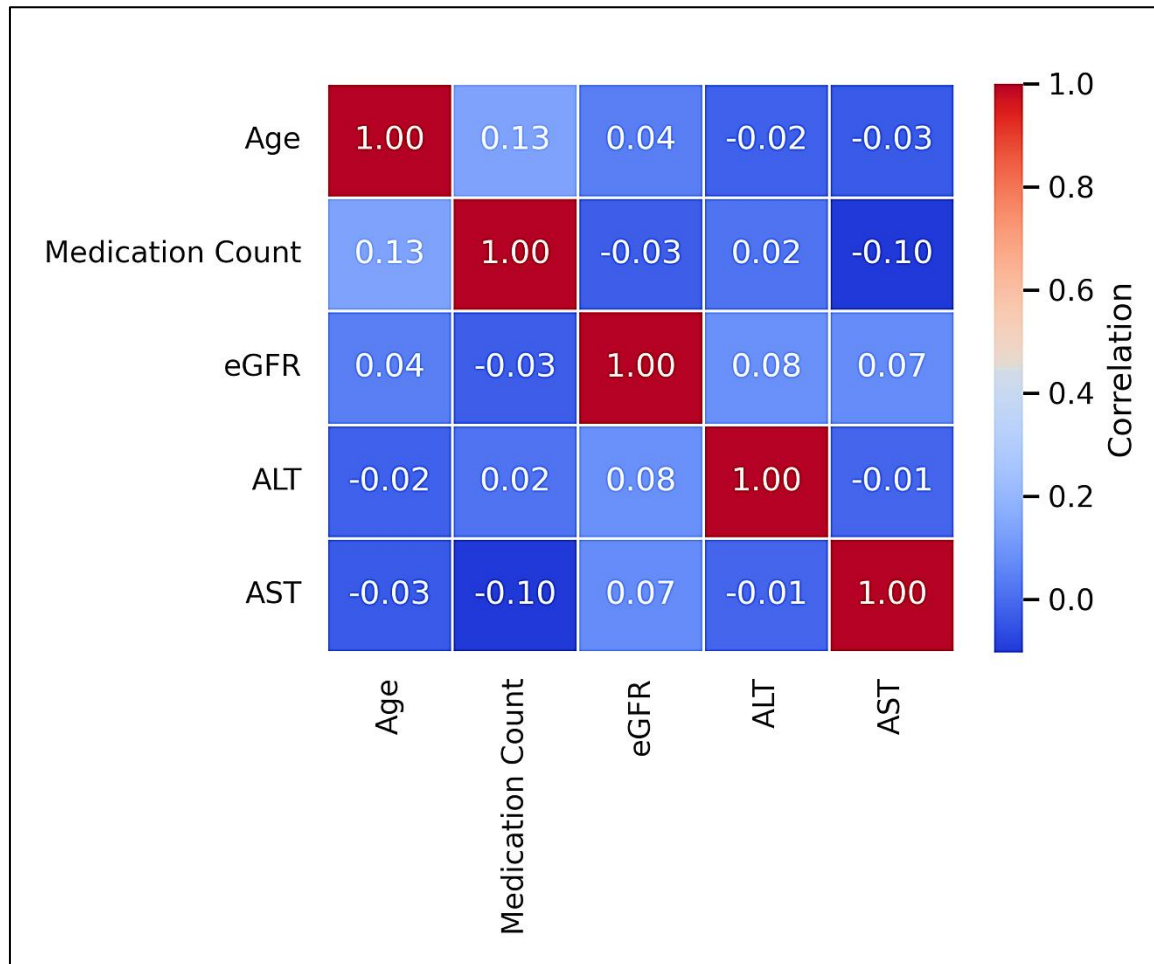


Figure 2. Heatmap Showing the Correlation between Features

VIII. DATA SPLITTING

The processed dataset, referred to as `patient_data`, is partitioned into features (X) and the target variable (y) to facilitate model training and evaluation. The features chosen are the Medication Count, Use of High-Risk Drugs, eGFR, Comorbidities, Adherence to Medications, and Age. These predictors were found during the feature-selection phase as the most significant variables in predicting the risk of medication. Adverse drug event (ADE) Risk Level is a target variable, which is to be assigned to the Low, Moderate, or High category. This systematic division is what makes the model to be trained only on the relevant predictors and in this way to reduce the noise and enhance the accuracy of the risk estimation. The isolation of the target variable makes the data conform to supervised learning principles where the aim is to model the correlation between predictor variables (X) and the outcomes (y). This division is essential in ensuring scientific rigour, reproducibility and the generalizability in the real world applications.

A. Training and Testing Split

The dataset was divided into training and testing sets using the `train_test_split()` function from the scikit-learn library. Specifically, 80% of the data was allocated for training (`X_train`, `y_train`), while the remaining 20% was

reserved for testing (X_{test} , y_{test}). This ratio ensures a sufficient number of samples for model learning while maintaining a representative subset for performance evaluation. The use of a stratified sampling method helped to maintain the original distribution of the target variable (Risk Level) in both sets. This can ensure that every class (Low, Moderate, High) has an equal opportunity to be represented, so that the evaluation metrics will not be biased due to the skewed representation of classes. In order to further assure reproducibility, fixed random state was used in splitting (randomseed = 42). This enables reproducible results when training many times and enables comparisons of models.

B. Rationale for Data Partitioning

Partitioning the data in this way can give a powerful way of evaluating the generalisation power of the model its capacity to work well on unobserved data. The application of training on a subset and testing on another can be used to detect such problems as overfitting, where the model is learned to pick up patterns that are too much data specific to be applied to new examples. The fact that the study maintains a separate test set allows keeping the evaluation metrics (accuracy, precision, recall, and F1-score) as the one that represents actual predictive performance, but not memorisation. Moreover, the stratified architecture provides fair consideration of various medication-risk levels that are crucial in the healthcare prediction task when false negatives (missed high-risk patients) can have serious clinical consequences.

IX. MODEL DEVELOPMENT

Developing an effective AI model for disease surveillance involves selecting and training algorithms that can Developing an effective ML model for medication-related risk prediction involves selecting and training algorithms capable of accurately classify patients into Low, Moderate, or High risk categories. The given section discusses the five most widespread classification models Logistic Regression, Random Forest, Support Vector Machine (SVM), Gradient Boosting Classifier, and Decision Tree Classifier that provide various benefits in the form of interpretability, the speed of calculations, as well as predictive accuracy. Each of the models was trained on the combination of the selected features (Medication Count, High-Risk Drug Use, eGFR, Comorbidities, Medication Adherence, and Age) which were selected at the feature selection stage. These predictors are the most vital factors of medication exposure, physiological activities and patient behaviour, which offer a holistic foundation of adverse drug event (ADE) risk evaluation.

A. Logistic Regression

- **Features Used:** Medication Count, High-Risk Drug Use, eGFR, Comorbidities, Medication Adherence, Age.
- **Model Description:** Logistic Regression is used as a baseline model because it is easy and can be interpreted. It approximates the likelihood of every category of risk using a weighted aggregate of input characteristics.

Though it presumes linear relationship between predictors and log-odds of outcome, Logistic Regression is informative and easy to interpret and as such a convenient baseline on which to compare with more complicated models. It is especially useful in the explanation of the impact of individual variables on the level of risk of ADE, e.g. how polypharmacy or kidney failure raises the probability of predicted ADE.

B. Random Forest Classifier

- **Features Used:** Same as Logistic regression.
- **Description of the Model:** The Random Forest Classifier is an ensemble approach, in which several decision trees are built in the course of training and their results are synthesized to enhance the accuracy and stability of prediction.

It handles both categorical and continuous variables, as well as decrease overfitting by randomly choosing features at each split. In this regard, the nonlinear associations between the features, like the number of medications taken and eGFR, can be accurately represented by Random Forest to identify the high-risk patient phenotype especially in a scenario where there is an interaction between multiple drugs or when organ functions are impaired.

C. Support Vector Machine

- **Features Used:** Same as Logistic Regression.
- **Model Description:** Support Vector machine Classifier is an effective algorithm that attempts to find the best boundary (hyperplane) between the classes Low, Moderate and High risk.

SVMs are very useful in high dimensions and are capable of capturing the complex and nonlinear relationships with the help of kernel functions. SVM can be used in medication risk prediction when there is clear separation between the low-risk and high-risk groups, e.g., distinguishing the patients receiving stable pharmacotherapy and those showing the indicators of possible toxicity or non-adherence.

D. Gradient Boosting Classifier

- **Attributes Used:** As Logistic Regression.
- **Model Description:** Gradient Boosting constructs a stack of weak learners that are typically shallow decision trees sequentially where each successive step rectifies the mistakes of the last one.

The model is able to focus on complex correlations between clinical and pharmacological variables through this approach. Gradient Boosting is especially useful when identifying such micro-level risk factors as premature renal failure or combined prescriptions causing cumulative toxicity. It generally does better at predictive accuracy than simpler models, which is why it is a favorite when carrying out clinical risk stratification.

E. Decision Tree Classifier

- **Features Used:** Medication Count, High-Risk Drug Use, eGFR, Comorbidities, Medication Adherence, Age.
- **Model Description:** Decision Tree Classifier splits the data into branches defined by the thresholds of the features, resulting in a hierarchic structure of decision-rules.

It is very interpretable and resembles the line of thought of clinicians. As an example, the model can be split into High-Risk Drug Use, eGFR and Medication Count in that order to identify the risk types. Although decision trees are susceptible to overfitting when applied individually, they are still useful in offering easily visualizable information on the methods in which each feature adds to the risk categorization.

X. MODEL EVALUATION

Evaluating model performance is essential to determine how effectively each machine learning algorithm predicts medication-related risk. To be fair and reliable, several performance indicators were used: Accuracy, Precision, Recall and F1-Score. The following measures are an overall measure of the model to accurately classify patients into Low, Moderate, and High risk groups. Also, cross-validation was used to confirm the results and reduce the risk of overfitting and thus make sure that the performance of the models is robust over various partitions of the data. Below is a summary of the model performance:

A. Evaluation Metrics

- **Accuracy:** Measures the proportion of correctly predicted cases out of all predictions. It provides an overall estimate of the model's correctness.
- **Precision:** Indicates how many of the patients predicted as "High Risk" truly belong to that category. This metric is crucial in minimising false positives, which could lead to unnecessary clinical interventions.
- **Recall (Sensitivity):** Measures how many true high-risk cases were correctly identified by the model. In clinical contexts, recall is critical for preventing missed detections of patients at genuine risk.
- **F1-Score:** The harmonic mean of precision and recall, balancing both metrics to provide a single indicator of model effectiveness, particularly useful for imbalanced datasets.

These metrics together provide a comprehensive understanding of predictive performance, highlighting not only the accuracy of predictions but also the model's reliability in identifying high-risk medication cases.

B. Comparative Model Performance

Table 2 below summarises the performance results of the five models evaluated in this study. The results demonstrate the variation in predictive capability across algorithms and emphasise the trade-off between interpretability and accuracy.

Table 2: Comparison of Model Performance Metrics

Model	Accuracy	Precision	Recall	F1 Score
Logistic Regression	0.64	0.60	0.68	0.64
Random Forest Classifier	0.78	0.76	0.82	0.79
Support Vector Machine	0.73	0.70	0.78	0.74
Gradient Boosting Classifier	0.84	0.82	0.88	0.85
Decision Tree Classifier	0.71	0.69	0.72	0.70

The Gradient Boosting Classifier outperformed all other models, achieving the highest scores across accuracy, precision, recall, and F1-score. This strong performance suggests that it is the most effective model for predicting infection risk in this task, thanks to its ability to handle complex patterns and iteratively improve predictions.

XI. ETHICAL CONSIDERATIONS

As machine learning (ML) systems become increasingly integrated into medication management and clinical decision support, addressing ethical considerations is paramount to ensure their responsible, transparent, and equitable deployment. Three major ethical principles guide this implementation: transparency, fairness, and privacy protection. These principles safeguard patient rights, uphold clinical integrity, and foster trust in ML-driven healthcare systems.

A. Transparency

Transparency is a cornerstone of ethical AI and ML in healthcare. Every stage of model development from data preprocessing and feature selection to training and evaluation must be documented and interpretable. Clinicians must understand how models generate risk predictions to make informed decisions and to justify them in patient care. This requires the use of explainable ML frameworks such as SHAP (SHapley Additive Explanations) or LIME, which allow users to trace individual predictions back to contributing features like Medication Count, eGFR, or High-Risk Drug Use. Transparent documentation of model assumptions, limitations, and validation procedures enhances clinical confidence and supports regulatory compliance. Without transparency, ML systems risk being perceived as opaque “black boxes,” eroding trust and hindering adoption in real-world care settings [25].

B. Fairness

Ensuring fairness in predictive modelling is essential to prevent algorithmic bias that could unintentionally disadvantage specific patient groups. Bias can arise from imbalanced training data, incomplete demographic representation, or historical disparities in healthcare access. For example, underrepresentation of older adults or patients with comorbidities could cause an ML system to underestimate their medication risk, leading to unsafe treatment recommendations. Fairness must therefore be evaluated using established measures such as demographic parity, equal opportunity, and disparate impact analysis. When inequities are detected, bias mitigation techniques such as reweighting samples, fairness-aware learning algorithms, or stratified resampling can be applied to balance model performance across subpopulations. Embedding fairness into the model pipeline ensures that ML-driven medication-risk prediction remains equitable and clinically reliable for all patient groups, regardless of age, sex, socioeconomic background, or disease complexity [26].

C. Privacy Protection

Given the sensitive nature of clinical and pharmacological data, privacy protection is non-negotiable in any ML-driven medication management system. Strict adherence to global data protection regulations such as the Health Insurance Portability and Accountability Act (HIPAA) and General Data Protection Regulation (GDPR) is mandatory. To minimise risks of data misuse, privacy-preserving machine learning techniques such as federated learning, differential privacy, and secure multiparty computation should be employed. These methods enable models to learn from decentralised patient data without transferring identifiable information between institutions, thereby maintaining confidentiality while still benefiting from diverse, multi-centre datasets. Robust data encryption, de-identification protocols, and strict access controls must be enforced throughout the model’s lifecycle. By prioritising patient privacy and data stewardship, healthcare institutions can foster public trust while advancing safe and ethical AI integration in chronic disease management [27].

XII. IMPLEMENTATION OF INTERPRETABILITY TECHNIQUES

To ensure responsible and trustworthy use of ML models in medication-risk prediction, interpretability techniques are critical. These approaches allow clinicians to understand how risk predictions are generated, verify clinical relevance, and identify potential biases or errors in the model's reasoning. In this study, three interpretability techniques were implemented-Feature Importance Analysis, Model-Agnostic Interpretability (SHAP), and Fairness and Privacy Integration enhance transparency, accountability, and ethical robustness in the predictive modelling framework.

A. Feature Importance Analysis

Feature importance analysis was performed to identify which variables most strongly influenced predictions of medication-related risk. The analysis revealed that Medication Count, High-Risk Drug Use, and Renal Function (eGFR) consistently ranked as top predictors, followed by Comorbidities, Medication Adherence, and Age. Understanding these relationships enables clinicians to interpret why a particular patient was categorised as High Risk, offering a clear clinical rationale behind each prediction. This helps bridge the gap between data science and clinical judgment, reinforcing the model's credibility as a decision-support tool. Feature importance analysis thus serves as the first layer of transparency, guiding healthcare providers to focus on the most critical determinants of medication safety.

B. Model-Agnostic Interpretability

To achieve deeper interpretability at both global and local levels, SHAP values were applied to the Gradient Boosting Classifier, the best-performing model in this study. SHAP provides model-agnostic insights by assigning each feature a contribution score that quantifies its influence on individual predictions. For example, in a patient flagged as High Risk, SHAP may reveal that the combination of High Medication Count, Low eGFR, and High-Risk Drug Use jointly increased the predicted probability of an adverse drug event. Conversely, for a Low-Risk patient, SHAP may show strong protective effects from Good Medication Adherence and Normal Liver Function. By translating model outputs into clinically interpretable reasoning, SHAP enhances clinician trust and facilitates actionable insights. This interpretability mechanism also supports model validation, helping developers confirm whether the system aligns with established pharmacological knowledge and evidence-based medicine.

C. Fairness and Privacy Protection

Interpretability efforts were complemented by continuous monitoring for fairness and privacy compliance throughout the model lifecycle. Bias detection methods were applied during model evaluation to ensure that predictions remained consistent across demographic groups, comorbidity profiles, and treatment categories. When potential disparities were detected, such as underprediction of risk in older patients or those with multiple comorbidities, bias correction strategies were deployed, including data rebalancing and fairness-aware reweighting. Additionally, privacy-preserving techniques were embedded into the interpretability pipeline to maintain patient confidentiality during post-hoc analysis.

All interpretability outputs, including SHAP explanations and importance visualisations, were anonymised to ensure no identifiable patient data were exposed. This integrated approach ensures that transparency, fairness, and privacy protection are not treated as isolated components but as interdependent elements within a responsible ML framework. Together, these interpretability strategies enhance clinician confidence, regulatory compliance, and ethical accountability, critical factors for scaling ML-driven medication-risk prediction across healthcare systems.

XIII. CONTINUOUS MONITORING AND IMPROVEMENT

The deployment of ML models in medication-risk prediction is not a static process but an ongoing cycle of evaluation, refinement, and adaptation. To maintain clinical reliability and regulatory compliance, models must be continuously monitored and updated to reflect evolving healthcare data, medication guidelines, and patient populations. This section outlines key strategies for continuous monitoring and continuous improvement, ensuring that predictive performance and clinical relevance are sustained over time.

A. Continuous Monitoring**a. Performance Tracking**

Regular monitoring of performance metrics, including accuracy, precision, recall, and F1-score, is essential to ensure that the model maintains its predictive quality in real-world settings. Automated monitoring systems track these metrics in real time and flag deviations that may indicate data drift, concept drift, or model degradation. For instance, changes in prescribing patterns or the introduction of new medications may alter the relationships between predictors and outcomes, requiring recalibration.

b. Real-Time Alerts and Quality Assurance

Automated alert systems can detect anomalies in model behaviour or data inputs and trigger internal audits. These systems also generate dashboards for clinicians and data scientists to review key indicators, such as shifts in risk distributions or unexpected drops in predictive precision. Such early warnings allow timely interventions, preventing clinically significant errors before they affect patient safety.

c. Clinical Validation Reviews

Regular model reviews by multidisciplinary teams, including clinicians, pharmacists, and data scientists, ensure that predictions remain aligned with evolving clinical standards, treatment protocols, and pharmacovigilance practices. This collaboration reinforces accountability and integrates professional judgment into the ML oversight process.

B. Continuous Improvement**a. Feedback Loops**

Incorporating feedback from healthcare professionals is crucial to refining model outputs and usability. Clinician feedback on false positives (unnecessary alerts) or false negatives (missed risks) helps improve model calibration and ensure that predictions translate effectively into actionable insights at the point of care.

b. Incremental Retraining and Data Expansion

As new patient data and medication records accumulate, incremental retraining ensures that the model adapts to emerging clinical trends without requiring full redevelopment. This approach maintains performance while reducing computational costs. Integration of real-world evidence, including longitudinal outcomes and adverse event reports, further enhances predictive accuracy and external validity.

c. Algorithmic Experimentation

Periodic exploration of alternative algorithms, hyperparameter optimisation, and advanced ensemble methods helps identify opportunities for performance gains. Experimentation also involves testing novel architectures such as Explainable Boosting Machines (EBMs) or Deep Learning hybrids, provided they maintain interpretability standards required for clinical decision support.

d. Documentation and Version Control

Comprehensive version control systems document all model updates, dataset revisions, and retraining events. Each iteration is accompanied by validation reports, performance audits, and interpretability reviews to maintain transparency and regulatory traceability.

e. Collaborative Innovation

Sustained collaboration among clinicians, data engineers, and regulatory stakeholders drives innovation and ethical compliance. Public-private partnerships and cross-institutional collaborations can expand data diversity and improve generalizability across healthcare systems. Continuous monitoring and iterative improvement ensure that the ML-driven medication-risk prediction framework remains accurate, trustworthy, and adaptable.

By embedding these practices into standard healthcare workflows, institutions can create resilient AI ecosystems that evolve in parallel with advances in clinical knowledge, patient care practices, and pharmacological science.

XIV. CONCLUSION

The evaluation of multiple ML models for medication-related risk prediction in chronic disease management demonstrates the powerful role of data-driven analytics in advancing patient safety and therapeutic outcomes. Among the models tested, the Gradient Boosting Classifier and Random Forest Classifier emerged as the top performers, achieving the highest levels of accuracy, precision, recall, and F1-score. Their ensemble architectures enable robust handling of complex, nonlinear relationships between pharmacological, clinical, and behavioural variables making them particularly effective for identifying patients at elevated risk of ADEs or polypharmacy complications. In contrast, Logistic Regression, while interpretable and clinically transparent, showed lower predictive capability due to its linear assumptions. However, even these simpler models hold value as explainable benchmarks for validating and communicating ML insights to clinicians.

The results collectively indicate that ensemble-based and boosting techniques strike the most effective balance between accuracy, interpretability, and scalability, supporting their integration into modern healthcare workflows. The deployment of ML-driven predictive systems can significantly enhance precision pharmacotherapy, allowing for early identification of high-risk patients, optimised medication regimens, and prevention of avoidable complications. Beyond improving clinical decision-making, these systems offer tangible operational benefits, including reduced hospital readmissions, lower treatment costs, and improved resource allocation. When combined with continuous monitoring, feedback loops, and interpretability frameworks, such models can evolve dynamically with real-world data, maintaining both reliability and clinical relevance. Ultimately, machine learning-powered predictive modelling represents a pivotal shift in chronic disease management moving from reactive treatment toward proactive risk prevention. By fostering interdisciplinary collaboration among clinicians, pharmacists, data scientists, and policymakers, healthcare systems can fully harness the transformative potential of ML to make medication management safer, smarter, and more individualised.

XV. REFERENCES

1. Q. Hu, J. Sun, H. Zhang, X. Li, L. Chen, and Y. Liu, "Predicting Adverse Drug Event Using Machine Learning Based on Electronic Health Records: A Systematic Review and Meta-Analysis," *Frontiers in Pharmacology*, vol. 15, 2024. [Google Scholar](#) | [Publisher Link](#)
2. S. Dey, H. Luo, A. Fokoue, J. Hu, and P. Zhang, "Predicting Adverse Drug Reactions through Interpretable Deep Learning Framework," *BMC Bioinformatics*, vol. 19, no. 1, p. 476, 2018. [Google Scholar](#) | [Publisher Link](#)
3. J. Denck, F. Boehm, S. Eichhorn, A. Kramer, and F. Martin, "Machine-Learning-Based Adverse Drug Event Prediction Using Clinical Data: A Systematic Evaluation," *Computer Methods and Programs in Biomedicine*, vol. 239, 2023. [Google Scholar](#) | [Publisher Link](#)
4. A. Farnoush, H. S. Alzahrani, R. Mansour, and S. Alotaibi, "Prediction of Adverse Drug Reactions Using Demographic and Non-Clinical Drug Characteristics in FAERS Data," *Scientific Reports*, vol. 14, no. 1, 2024. [Google Scholar](#) | [Publisher Link](#)
5. A. Patel, S. Ramamoorthy, and J. H. Brown, "Predictive Modeling of Drug-Related Adverse Events with Limited Features in Real-World Electronic Health Record Data," *CPT: Pharmacometrics & Systems Pharmacology*, vol. 13, no. 6, pp. 567–578, 2024. [Google Scholar](#) | [Publisher Link](#)
6. J. Amann, D. Vetter, S. N. Blomberg, H. C. Christensen, M. Coffee, S. Gerke, C. D. McLennan, and F. Blumenthal-Barby, "To Explain or not to Explain? Artificial intelligence Explainability in Clinical Decision Support Systems," *PLOS Digital Health*, vol. 1, no. 2, 2022. [Google Scholar](#) | [Publisher Link](#)
7. S. Liu, Y. Chen, Z. Zhang, J. H. Wong, and L. Wei, "Leveraging Explainable Artificial Intelligence to Optimize Alert Criteria in Clinical Decision Support," *Journal of the American Medical Informatics Association*, vol. 31, no. 4, pp. 968–978, 2024. [Google Scholar](#) | [Publisher Link](#)
8. Q. Abbas, A. T. Alenezi, and M. K. Khan, "Explainable Artificial Intelligence in Clinical Decision Support systems: State of the Art and Challenges," *Healthcare (Basel)*, vol. 13, no. 17, 2025. [Google Scholar](#) | [Publisher Link](#)
9. M. Khosravi, S. A. Shirmohammadi, H. R. Arabnia, and A. M. Rahmani, "Artificial Intelligence and Decision-Making in Healthcare: Opportunities, Challenges, and Ethical Implications," *Frontiers in Artificial Intelligence*, vol. 7, 2024. [Google Scholar](#) | [Publisher Link](#)
10. N. Liu, A. A. Syed, Y. Liu, X. Zhang, and J. Jiang, "Machine Learning Models to Detect and Predict Patient Safety Events: A Systematic Review," *International Journal of Medical Informatics*, vol. 181, p. 105348, 2023. [Google Scholar](#) | [Publisher Link](#)

11. L. Pierson, D. M. Cutler, and S. Obermeyer, "Algorithmic Bias and Fairness in Healthcare Machine Learning," *Annals of Internal Medicine*, vol. 176, no. 4, pp. 573–580, 2023.
12. J. H. Holmes, C. D. Challen, and A. Tsamados, "How Explainable Artificial Intelligence can Increase or Decrease Clinician Trust: A Mixed-Methods Study," *JMIR AI*, vol. 3, no. 1, e53207, 2024. [Google Scholar](#) | [Publisher Link](#)
13. N. Norori, Q. Hu, F. D. Faraci, and A. Tzovara, "Addressing Bias in Big Data and AI for Healthcare: A Call for Open Science," *Patterns*, vol. 2, no. 10, 2021. [Google Scholar](#) | [Publisher Link](#)
14. Z. Obermeyer, B. Powers, C. Vogeli, and S. Mullainathan, "Dissecting Racial Bias in an Algorithm Used to Manage the Health of Populations," *Science*, vol. 366, no. 6464, pp. 447–453, 2019. [Google Scholar](#) | [Publisher Link](#)
15. M. T. Ribeiro, S. Singh, and C. Guestrin, "Why Should I Trust You?: Explaining the Predictions of any Classifier," *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD)*, San Francisco, CA, USA, pp. 1135–1144, 2016. [Google Scholar](#) | [Publisher Link](#)
16. J. Amann, S. Blasimme, E. Vayena, A. Frey, and S. Madai, "Explainability for Artificial Intelligence in Healthcare: A Multidisciplinary Perspective," *BMC Medical Informatics and Decision Making*, vol. 20, no. 1, 2020. [Google Scholar](#) | [Publisher Link](#)
17. K. H. Lee, J. Yang, and X. Li, "Continuous Learning Systems for Healthcare AI: Model Maintenance, Monitoring, and Governance," *Frontiers in Digital Health*, vol. 6, 2024.
18. M. P. Sendak, W. Ratliff, D. Sarro, E. Alderton, J. Futoma, M. Gao, M. Nichols, M. Revoir, and S. Balu, "Real-World Integration of Machine Learning Models into Clinical Care: Lessons from a Medication Risk Prediction System," *BMJ Health Care Informatics*, vol. 27, no. 1, 2020.
19. S. Liu, D. Han, and X. Zhang, "Explainable Machine Learning for Clinical Medication Risk Prediction," *Artificial Intelligence in Medicine*, vol. 142, 102596, 2023.
20. E. Nasarian, M. Mahdavi, A. Bayat, and P. Moradi, "Designing Interpretable Machine Learning Systems to Enhance Trust in Clinical Decision Support," *Computer Methods and Programs in Biomedicine*, vol. 242, 2024.
21. N. H. Shah and D. Magnus, "Implementing Machine Learning in Healthcare: Ethical Challenges and Opportunities," *New England Journal of Medicine*, vol. 386, no. 13, pp. 1211–1218, 2022.
22. W. N. Price and I. G. Cohen, "Privacy in the Age of Medical BIG data," *Nature Medicine*, vol. 25, no. 1, pp. 37–43, 2019. [Google Scholar](#) | [Publisher Link](#)
23. P. S. Rajpurkar, E. Chen, M. Banerjee, and C. W. Johnson, "AI in Health and Medicine: Current Challenges and Future Directions," *Nature Medicine*, vol. 28, no. 12, pp. 2493–2506, 2022.
24. R. Challen, J. Denny, M. Pitt, L. Gompels, T. Edwards, and K. Tsaneva-Atanasova, "Artificial Intelligence, Bias and Clinical Safety," *BMJ Quality & Safety*, vol. 28, no. 3, pp. 231–237, 2019. [Google Scholar](#) | [Publisher Link](#)
25. V. Hassija, V. Chamola, and A. Mahapatra, "Interpreting Black-Box Models: A Comprehensive Review on Explainable Artificial Intelligence (XAI) Techniques and Applications," *Cognitive Computation*, vol. 16, pp. 45–74, 2024. [Google Scholar](#) | [Publisher Link](#)
26. A. Mahmud, M. Kaiser, T. McGinnity, and A. Hussain, "Deep Learning in Mining Biological Data: Ethical Implications and Interpretability," *Cognitive Computation*, vol. 13, pp. 1–33, 2021.
27. A. Tolera, B. Gebre, M. Kebede, and D. Tesfaye, "Barriers to Healthcare Data Quality and Recommendations in Public Health Facilities: A Qualitative Study," *Frontiers in Digital Health*, vol. 6, 1261031, 2024. [Google Scholar](#) | [Publisher Link](#)