

Automated Software Bug Severity Classification: A Comparative Performance Analysis of ML and DL Models

Sweety Ahlawat¹, Dr. Dhiraj Khurana², Dr. Yogesh Kumar³

¹ Department of Computer Science and Engineering, Maharshi Dayanand University, Rohtak – 124001, Haryana, India.
Email: sweetyahlawat6@gmail.com

² Department of Computer Science and Engineering, Maharshi Dayanand University, Rohtak – 124001, Haryana, India.
Email: dhirajkhurana@mdurohtak.ac.in

³ Department of Computer Science and Engineering, Maharshi Dayanand University, Rohtak – 124001, Haryana, India.
Email: dryogeshkumar.uiet@mdurohtak.ac.in

Abstract: In order to improve software quality, optimize resources, and minimize system failure, it is crucial to properly classify the severity of software issues. A comprehensive plan for automated bug severity classification using a weighted combination of source code metrics, class imbalance handling using the SMOTE algorithm, and a comparison of ML and DL models is proposed in this paper. A number of models were evaluated using the usual performance metrics, such as ROC-AUC, F1-score, recall, accuracy, and SVM, CNN, and LSTM. The experimental results demonstrate that SMOTE significantly improves recall and F1-score by increasing the detection of minority classes, particularly important faults. The capacity to reflect more complex trends in software measurements was proved by LSTM, which outperformed all other models with an accuracy of 89.2% and F1-score of 0.86. Research also provides further data by means of stability, robustness, and class-wise analyses, which demonstrate that DL models are better than conventional ML models. The suggested solution might make the problem triage process easier, better, and more automatic in real-world development settings. This might lead to faster debugging, better choices, and more dependable goods.

Keywords: Software Bug Classification, Bug Severity Prediction, Machine Learning, Deep Learning, SMOTE, Data Imbalance, Automated Bug Triage

1. Introduction

Software quality assurance is a very important part of current software engineering in which detection and prioritization of defects in software at an early stage affect a lot the reliability of the software system as well as the software maintenance cost. Bug severity classification is among many other defect management activities that are critical in establishing urgency and impact of reported bug [1]. Proper severity forecasting helps a developer to manage the resources and respond to the critical defects in time. Conventionally, the evaluation of bug severity is done manually whereby the assessment process is subjective, time consuming and also subjected to inconsistency. The main attribute of such predictive systems is modeling software artifacts in terms of source code metrics, binding quantitative measures of the code complexity, size and structure [2]. LOC, cyclomatic complexity and Halstead measurements are the metrics that reflect key attributes of the software and are used as informative measures to the classification models. Nevertheless, class imbalance, i.e. the lack of certain areas of severity levels (e.g., critical bugs) is one of the largest problems with bug severity classification. This imbalance may cause the predictive models to favor majority classes, creating poor generalization and poor performance on the minority classes. To solve this problem, SMOTE (a data-level method) is used to create synthetic data and normalize the distribution of data [3]. Although the application of ML and DL methods has increased, there is still a gap in extensive literature examining the performance of these

algorithms in the imbalanced data in terms of source code measures. Additionally, the effect of SMOTE on the various learning paradigms has not been well tested in a coherent experimental design. Thus, the paper is based on a comparative performance of SMOTE-augmented ML and DL models to automatically classify severity of software bugs [4]. This main aim is to determine how well SMOTE can enhance the classification performance and the best learning style that should be used in this case.

1.1 Background and motivation

As software systems have grown in scale and complexity, software reliability is becoming more difficult to assure. The recent applications have many bugs to be reported to at both stages of development and after it is deployed and it is not easy to analyze them manually and prioritize them properly as expected by the developers [5]. Bug severity classification is one of many activities in defect management that are necessary to figure out what issues are most critical and need to be taken seriously and consequently directly affect the quality of the software and customer satisfaction. Conventionally, the allocation of the severity of bugs is done by human specialists according to his/her experience and knowledge of the system [6]. This process is however more often than not subjective, inconsistent and cannot be scaled to large scale software projects. Consequently, increasing pressure is on automated methods that have the ability to rank the severity of bugs correctly and reliably. ML and DL techniques have come out as potential solutions to this issue because, on historical data, they can learn complex patterns. One of the basic conditions of the implementation of these methods is the conversion of raw software products into structured representations. This is done by measuring source code metrics, which is a measure of different attributes of the code including its size, complexity and maintainability [7]. Such measures as LOC, cyclomatic complexity, and Halstead measures offer useful information on the fault-proneness of server software modules and can be used as effective predictive modeling attributes. Even with these developments, class imbalance has become one of the most problematic issues in the classification of bugs where high severity bugs are greatly outnumbered by low severity bugs. Such imbalance may result into biased models that can be poor when applied to minority classes which are usually the most important. In order to solve this problem, synthetic minority oversampling SMOTE are used to create artificial samples and equalize the data [8]. The rationale of this study is based on the necessity to conduct a systematic assessment of the functioning of various paradigms of learning in such circumstances. Although some of the studies have evaluated ML and DL models separately, no attempt has been made to create a comparative analysis of these models that consider the measures of source codes, the management of class imbalances via SMOTE, and the performance of these models. It is important to understand these interactions so as to establish strong and stable automated systems through which the severity of bugs can be classified [9].

1.2 Importance of bug severity prediction in software engineering

Bug severity prediction is considered to be a core component of software quality assurance because it decides the effect and urgency of errors in a software system. Proper prediction of bug severity allows development teams to concentrate on the critical defects which can have a severe impact on the functionality of the system, its security or user experience. Efficient allocation of resources is one of the key factors that make bug severity prediction significant [10]. The software development teams usually have time and budget constraints and it is best to spot high severity bugs early so that the issues of significance are dealt with early and those of lesser concern can be left. This priority enhances production and minimizes the chances of failure of systems in manufacturing. The other important issue is the increased software quality and dependability. Bugs of high severity that are not addressed might result in crashing the system, loss of data or security breaches. Precise foresight of the degree of severity enables the organization to take forward-looking measures to debug its systems, thus, making them very stable and reducing risks to the minimum. Bug severity prediction is also a contributor to speedy decision making and less development time. Automated classification, especially on the basis of ML and DL does not require manual work, and it also gives consistent results [11]. Moreover, when working with huge and skewed datasets, methods like SMOTE are more effective towards detecting rare yet critical bugs. This guarantees that minority classes that are usually high severity defects are not neglected during model training. The other significant advantage is that it improves maintainability and lowers costs. Early detection of serious flaws will mean the cost of bug fixing will be low since it is normally detected at the late stages that are usually costly and complicated as well. It is also used in ensuring cleaner and more stable codebases in the long term. Lastly, the prediction of bugs assistance backs up software engineering that is based on data [12]. Predictive models are able to determine the pattern related to severe defects by utilizing the source code metrics of complexity, size and structural attributes. This not only increases the accuracy of classification, but also gives information on the quality of codes and possible areas of risk [13]. Table 1 identifies the major relevance of bug severity prediction to software engineering concepts, and it is used in prioritizing defects, resource allocation, and

minimization of risks. It describes the positivity of severity prediction in the enhancement of the quality of software, quick decision making, and cost reduction of the development. In general, it shows the relevance of automated and data-driven methods to the current software quality assurance.

Table 1: Importance of Bug Severity Prediction in Software Engineering

Aspect	Description	Impact on Software Engineering
Prioritization of Defects	Identifies and classifies bugs based on severity level	Helps teams focus on critical issues first
Resource Allocation	Enables efficient use of time, budget, and workforce	Improves productivity and cost management
Software Quality Improvement	Ensures high-severity bugs are resolved early	Enhances system reliability and performance
Risk Reduction	Prevents system crashes, data loss, and security breaches	Minimizes operational and security risks
Faster Decision-Making	Automates bug classification using ML/DL models	Supports quick and consistent decisions
Handling Imbalanced Data	Techniques like SMOTE improve detection of rare critical bugs	Ensures minority (high severity) bugs are not ignored
Reduced Development Time	Early detection reduces debugging time in later stages	Speeds up software development lifecycle
Cost Reduction	Fixing bugs early is cheaper than fixing in later stages	Lowens maintenance and operational costs
Improved Maintainability	Helps maintain cleaner and stable codebases	Supports long-term software sustainability
Data-Driven Engineering	Uses source code metrics to predict severity patterns	Enables intelligent and evidence-based decisions
Insight into Code Quality	Identifies relationships between code complexity and defects	Helps detect high-risk code areas

1.3 Role of ML and DL in software bug classification

ML and DL have made great contributions to the science of software engineering, allowing the automation of the complex process of software bugs classification, such as bug severity prediction. These methods can offer the use of data-driven techniques that substitute the conventional manual techniques, which have lower accuracy, consistency, and scalability [14]. The ML models are essential in the process of acquiring patterns of bugs historically and metrics related to source codes. These models are based on structured sets of features that are produced by software artifacts and are based on relationships between code characteristics (e.g., complexity, size, coupling) and the severity of defects. ML methods prove especially efficient in cases where the data is well-crafted and has a comparatively moderate size as it is interpretable and trains faster. ANN, CNN and RNN/LSTM models are capable of learning complex, non-linear relationships not necessarily seen by traditional ML models. DL models prove useful in particular when using large datasets or require other data sources like textual bug reports or sequential code changes [15]. End-to-end learning allows them to do feature engineering by hand, which reduces the necessity of manual feature engineering and thus are powerful in the automated classification task. One of the benefits of combining ML and DL in bug classification is that they can be adapted to unequal datasets, which is often a problem with bug severity prediction. Together with the methods such as SMOTE, both ML and DL models can learn minority classes more accurately, which enhance the recall and F1-score of important bug types. Moreover, the comparative performance can be made using both ML and DL models and this is critical in determining the most effective method in various

circumstances [16]. The interpretation of ML models can be often more interpretable and cheaper to compute, whereas the predictive performance of DL models can be higher since they are able to model complex patterns. The significance of this trade-off is to consider both paradigms in a single framework. Moreover, these strategies enable lifelong learning and development. This means that as new bug information is received, models can be retrained to accommodate changing software systems to maintain their performance over time. This renders ML and DL very appropriate to be incorporated in contemporary software engineering utilities like DevOps and continuous integration pipelines [17].

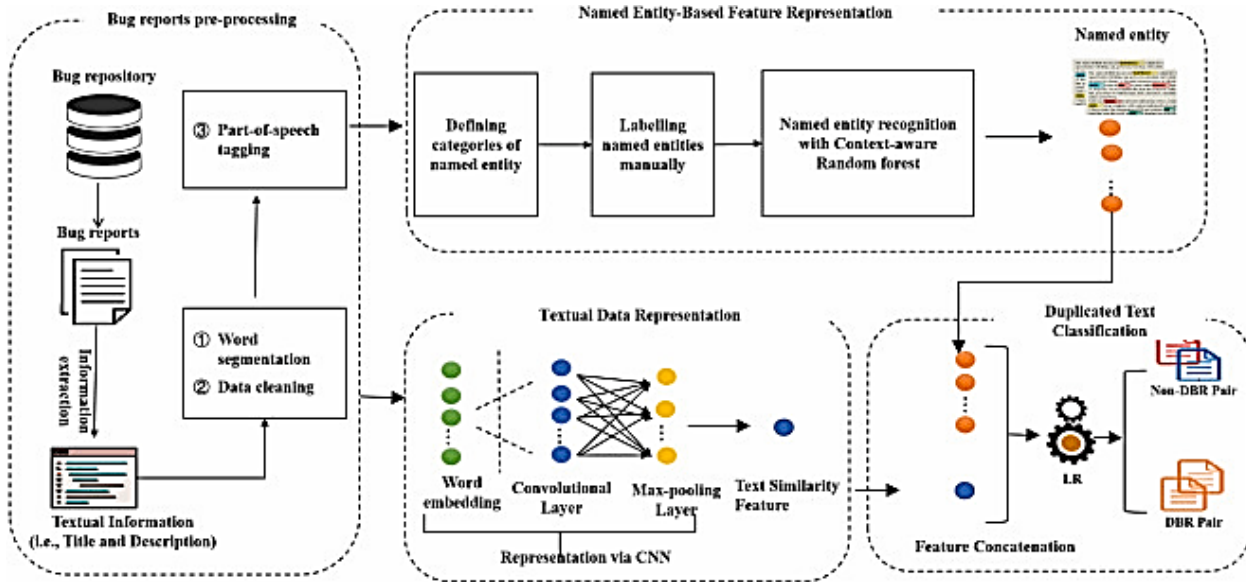


Figure 1: ML and DL in Software Bug Classification

As Figure 1 shows, ML and DL methods play an important role in automating software bugs classification. It points out that source code metrics and data on bugs are processed using various learning models to forecast the degree of severity. The figure also focuses on the workflow comparison of the ML and DL in finding patterns and enhancing classification accuracy. On the whole, it gives the conceptual summary of the intelligent bug triaging systems.

1.4 Challenges

Although there have been considerable improvements in automated software bug severity classification with the help of ML and DL, there are various issues that remain to be a barrier to the efficiency and quality of such solutions. Class imbalance, in which the dataset comprises low-severity bugs and a disproportionate number of high-severity or critical bugs, is one of the main problems, resulting in biased models that do not generalize to minority classes, despite such techniques as SMOTE being used [18]. The use of source code metrics to encode feature information also places constraints on the feature encoding since the metrics are only capable of capturing such properties of the structure as size and complexity, but fails to capture more semantic and contextual data, like developer intent and runtime behaviour. The other highly important issue is the quality of data because the bug repositories are often not complete, incorrectly labelled with the severity level, noisy, or redundant features, all of which adversely affect the model performance [19]. Moreover, the risk of choosing the right model is that there is a trade-off between complexity and generalization, on the one hand, deep learning model can fit the complex patterns, on the other, it is sensitive to overfitting and extremely expensive in terms of computation resources, whereas, on the other hand, traditional machine learning models might not be able to capture complex relationships. Moreover, even most advanced models, especially deep learning models, are not very interpretable, which are treated like black boxes and are challenging to justify predictions in very important software [20]. Lastly, generalization between software projects is typically problematic with models; differences between different coding standards, project areas, and project development practices can have a considerable impact on prediction quality.

1.5 Role of source code matrix in software bug classification

The source code metrics are essential in the automated software bug classification because they give an organized and quantitative view of software features. The metrics which are used to capture important features in the

code quality, size, structural complexity and related measures include LOC, cyclomatic complexity, Halstead measures, coupling, and cohesion, all of which are strongly linked to the occurrence and severity of defects [21]. These measures are the main input characteristics of classification models such that they can identify trends that make them learn to differentiate between various levels of bug severeness [22]. In addition, metrics on source codes can be used to implement data preprocessing processes like normalization and feature selection that enhances performance and stability of models. Regarding imbalanced datasets, these measures also constitute the feature space, over which such methods as SMOTE act to synthesize minority classes [23]. Moreover, the predictive models can be better interpreted by applying source code measures that enable the researcher and developers to gain insights into the connection between the variation of code features and severity of defects [24].

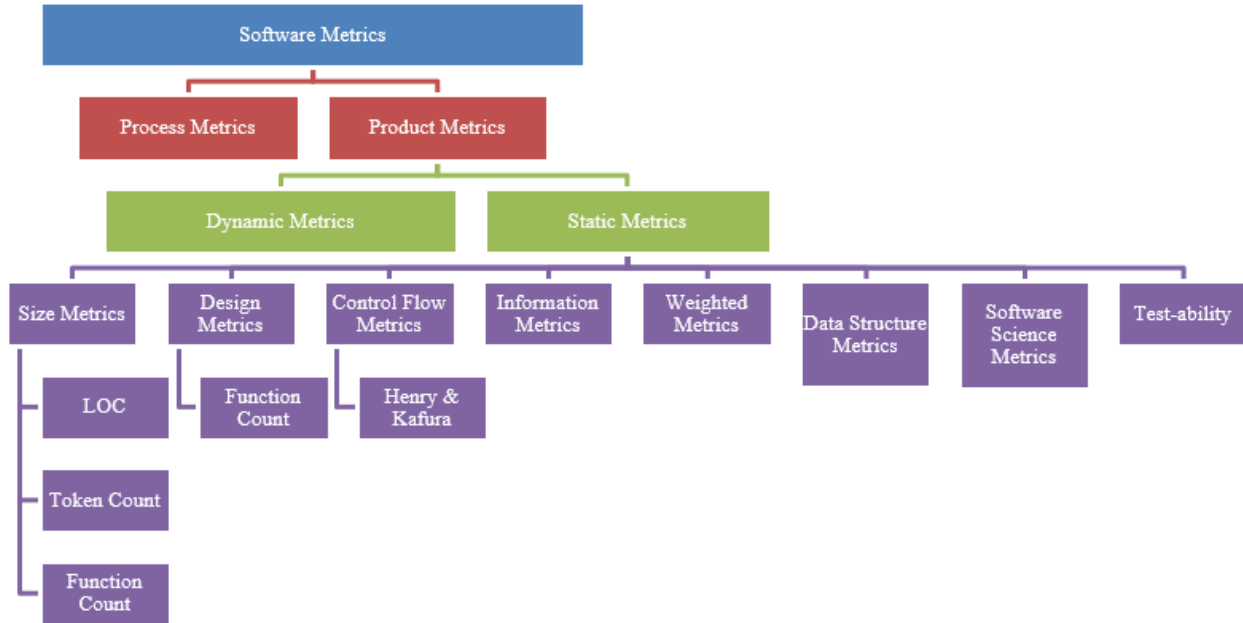


Figure 2: Source Code Matrix In Software Bug Classification

To classify software bugs into different categories, derived numerical features may be used to convert raw code into numerical features (see figure 2). It makes a distinction between important measures of code size, complexity, and structural properties including LOC, cyclomatic complexity, and Halstead measures. These characteristics are the very important inputs in ML and DL models to have a proper severity prediction. All in all, the figure highlights the relevance of representation of features in improving classification performance.

1.6 Contribution of research

The study provides an in-depth model of automated classification of source code bugs based on severity through the fusion of source code measures, class imbalance correction, and the combination of advanced learning models into a single experimental paradigm [25]. The main value of the research is that it provides a comparative analysis of performance of ML and DL models on a systematic basis of the feature space, which is based on source code metrics. This work presents a fair and structured comparison of the models unlike the other available studies, which test models individually, by using the same preprocessing methods and assessment guidelines on all of the models [26]. The use of SMOTE to overcome class imbalance issue and allow minority classes to be detected better is another meaningful contribution, especially those with high severity that can result in serious consequences in practical settings. The study also brings out the effect of feature representation by the aid of source code measures on classification and shows how the features improve the predictive accuracy and understandability [27]. Moreover, this piece of work offers the understanding of the trade-offs between the ML and DL models with respect to performance, computational complexity, and generalization ability [28]. On the whole, the presented framework will help to create better, more scalable, and more reliable bug severity prediction systems, which will facilitate the provision of data-driven decision-making during software quality assurance. The key novelty of this work lies in the unified evaluation of ML or DL models under SMOTE-augmented framework, complemented by advanced analytical insights such as stability, robustness, and bias–variance analysis [29]. The table 2 provides an overview of key contributions and contributions of new proposed research framework on automated bug severity classification. It emphasizes the incorporation of

source code measures, SMOTE-based class imbalance management, and the relative analysis of both ML and DL models. Besides, it highlights more sophisticated analytical knowledge including stability, robustness, and bias-variance analysis, which prove the general importance and novelty of the research.

Table 2: Key Contributions and Novelty of the Proposed Study

Aspect	Description	Research Value / Impact
Integrated Framework	Proposes a unified framework combining source code metrics, class imbalance handling (SMOTE), and ML/DL models	Provides a comprehensive and structured approach for bug severity classification
Comparative Analysis	Performs systematic comparison of ML and DL models under identical preprocessing and evaluation settings	Ensures fair, unbiased, and reliable performance evaluation
Class Imbalance Handling	Utilizes SMOTE to balance skewed datasets and improve minority class representation	Enhances detection of critical bugs, improving software reliability
Feature Representation	Uses source code metrics (LOC, complexity, Halstead) as input features	Improves model interpretability and predictive accuracy
Performance Improvement	Demonstrates significant gains in accuracy, recall, and F1-score after applying SMOTE	Leads to more effective and robust classification models
ML vs DL Trade-off Analysis	Evaluates models in terms of accuracy, computational complexity, and generalization	Provides practical insights for model selection
Advanced Analytical Insights	Includes stability analysis, robustness testing, and bias-variance evaluation	Strengthens scientific validity and depth of research
Generalization Capability	Assesses model consistency across different data splits and conditions	Ensures reliability in real-world applications
Scalability	Framework can be extended to large-scale datasets and real-time systems	Supports deployment in industrial environments
Practical Applicability	Enables automated bug triaging and data-driven decision-making	Reduces manual effort and improves software quality assurance
Novelty of Work	Unified evaluation of ML and DL models under SMOTE with advanced analysis techniques	Establishes a distinct contribution beyond existing studies

2. Literature Review

Fei et al. (2025) suggested a multivariate heterogeneous hybrid deep learning for predicting software faults [1]. Thakkar et al. (2025) created a self-explanatory machine learning to help control systems forecast defects in real time [2]. Cao (2025) developed a sophisticated network neural dynamics framework for automated bug triaging by using network-based representations to improve the effectiveness of bug categorization [3]. Yang et al. (2025) introduced a priority-sensitive LSTM with an attention mechanism to enhance bug priority prediction by integrating temporal dependency [4]. Dong et al. (2025) suggested a dynamic ensemble classifier selection that makes defect prediction more accurate [5]. Saraireh et al. (2025) developed a hybrid model using an adaptive ensemble learning and a binary white shark optimizer to improve the accuracy of software defect prediction categorization [6]. Anand et al. (2025) used wrapper-based dynamic arithmetic optimization for feature selection to enhance model efficiency and prediction accuracy [7]. Dharmakeerthi et al. (2025) showed that combining several deep learning could increase performance by employing a deep ensemble strategy to prioritize problems [8]. When Liu (2025) suggested a GAN-based SMOTE

oversampling method to fix class imbalance, synthetic samples made their ability to forecast faults better [9]. Mansouri et al. (2025) created a swarm-based cost-sensitive decision tree model to make classification rules better for datasets that aren't balanced and to make defect detection more accurate [10]. Quan et al. (2025) proposed XGV-BERT, which sought to make it easier to detect software system vulnerabilities by combining contextual language models with graph neural networks [11]. Tanveer et al. (2024) proposed a projection-based fuzzy LS-TSVM to enhance classification accuracy in contexts such as defect prediction [12]. Panda (2024) created a way to use intuitionistic fuzzy logic to guess how important errors are in order to deal with problems of ambiguity and class imbalance [13]. In a reexamination of code smell severity prioritizing by Liu et al. (2024) [14], ranking-based models enhanced the accuracy of prioritized. To enhance the predictive accuracy of software failure forecasting, Singh (2024) suggested the integration of novel code metrics with machine learning techniques [15]. Wang et al. (2024) conducted an empirical investigation on word embeddings and deep learning for bug assignment, illustrating the influence of representation schemes on performance [16]. Huang et al. (2024) explored how developer-related variables affect code smell assessment and showed how important human-centered factors are in prediction models [17]. Suryawanshi (2024) demonstrated a fundamental but efficient predictive methodology for software fault detection by logistic regression with gradient descent optimization [18]. Zhang et al. (2024) suggested a method to enhance the precision of defect predictions using RL [19]. This method uses deep Q-learning to get features. Fan et al. (2024) proposed a stable learning-based fault prediction model that focuses on making it more generalizable and resilient across diverse datasets [20].

To enhance the precision of fault prediction in intricate software systems, Qiao et al. (2024) presented DeMuVGN, a model based on GNN that collects data on software dependencies from several perspectives [21]. By integrating several feature representations, Sultan et al. (2024) improved classification effectiveness in multi-class software defect detection [22]. To better understand the context of text data, Bibyan et al. (2024) built a CNN-based model that used LDA and sentiment ratings from bug reports to estimate the severity of a problem [23]. For software defect datasets, Panda's (2024) intuitionistic fuzzy-based approach for bug seriousness prediction works well [24]. To forecast the severity of problems, Mian (2023) presented a CNN/BiLSTM hybrid model that finds sequential and geographic patterns in bug report data [25]. To improve the accuracy of classification, Dao (2023) developed an automated method for forecasting bug priority [26]. When it comes to unbalanced datasets, Hairani (2023) proposed a hybrid sampling approach that mixes SMOTE and ENN to improve classification accuracy [27]. By combining source code measurements with LLMs, Mashhadi et al. (2023) improved predictive accuracy via semantic understanding and created a model to predict the severity of issues at the method level [28]. By combining many classifiers, Johnson et al. (2023) improved software bug prediction using an ensemble learning approach [29]. By combining PCA feature selection with machine learning, Rao et al. (2023) improved model performance and addressed the issue of class imbalance in code smell severity detection [30]. Program semantic feature mining was proposed by Yao et al. (2023) as a fault identification method, with the goal of improving prediction accuracy by using deeper insights into the code level [31]. By comparing several machine learning for software defect prediction, Khalid et al. (2023) shed light on the advantages and disadvantages of various models [32]. Finding problems in imbalanced datasets became much simpler when Kaur (2023) developed SMOTE-based prediction model for cloud-aging software flaws [33]. Enhancing detection effectiveness in imbalanced cybersecurity datasets, Mahalakshmi et al. (2022) proposed a SCADA intrusion detection system that uses cost-sensitive machine learning and SMOTE-SVM [34]. Taking a look at granular defect patterns to improve prediction granularity, Mo et al. (2022) conducted an exploratory study into method-level bug prediction [35]. The effectiveness of text mining approaches in software defect analysis was shown by Hirsch (2022), who utilized textual bug reports to anticipate fault kinds [36]. Li et al. (2021) improved prediction effectiveness across varied datasets by introducing a cross-project bug prediction approach that uses joint probability domain adaptation and k-means SMOTE [37]. The effectiveness of transformer-based models in software maintenance operations was shown by Mashhadi (2021) who used CodeBERT to automatically fix fundamental Java mistakes [38]. Using information retrieval and machine learning, Neysiani (2020) investigated ways to detect duplicate bug reports, demonstrating the advantages of combining the two kinds of algorithms [39]. An artificial immune network hyper-parameter optimization method was proposed by Khan et al. (2020) to enhance the classifier's performance in software bug detection [40].

Table 1 provides a comparative review of available studies in software bug severity and defect prediction, providing the major approaches, methodologies, and results. It provides an overview of the latest innovations in ML, DL, and SMOTE-based methods, their advantages and disadvantages. The given comparison offers a solid background of the recognition of the gaps in the research and the explanation of why the given research should be undertaken.

Table 3: Comparison of Existing Researches

Ref	Author / Year	Objective	Methodology	Key Findings	Limitation
[1]	Fei et al. (2025)	Improve defect prediction accuracy	Hybrid deep learning (multivariate heterogeneous)	High accuracy with diverse features	High computational complexity
[2]	Thakkar et al. (2025)	Real-time defect prediction with explainability	Explainable ML framework	Improved transparency and decision-making	Limited scalability in large systems
[3]	Cao & Cui (2025)	Automate bug triaging	Complex network neural dynamics	Efficient bug classification	Complex model design
[4]	Yang et al. (2025)	Bug priority prediction	LSTM + Attention	Better contextual learning	Requires large training data
[5]	Dong et al. (2025)	Optimize classifier selection	Ensemble learning	Improved prediction performance	Model selection overhead
[6]	Saraireh et al. (2025)	Enhance defect classification	Ensemble + white shark optimizer	High classification accuracy	Optimization complexity
[7]	Anand et al. (2025)	Feature selection improvement	Wrapper-based optimization	Better feature relevance	Increased computation time
[8]	Dharmakeerthi et al. (2025)	Bug priority prediction	Deep ensemble models	Improved robustness	Resource-intensive
[9]	Liu & Liu (2025)	Handle class imbalance	GAN-based SMOTE	Better synthetic data quality	Training instability
[10]	Mansouri et al. (2025)	Improve imbalanced classification	Swarm-based cost-sensitive tree	Better minority class detection	Model complexity
[11]	Quan et al. (2025)	Vulnerability detection	BERT + GNN	Strong contextual understanding	High computational cost
[12]	Tanveer et al. (2024)	Address class imbalance	Fuzzy LS-TSVM	Improved classification	Limited generalization
[13]	Panda (2024)	Bug priority prediction	Intuitionistic fuzzy model	Handles uncertainty well	Complex parameter tuning
[14]	Liu et al. (2024)	Code smell prioritization	Learning-to-rank	Better prioritization accuracy	Data dependency
[15]	Singh (2024)	Improve fault prediction	ML + new metrics	Enhanced prediction accuracy	Feature engineering needed
[16]	Wang et al. (2024)	Bug assignment analysis	DL + word embeddings	Better representation learning	Model sensitivity
[17]	Huang et al. (2024)	Code smell prioritization	Developer feature analysis	Importance of human factors	Limited automation
[18]	Suryawanshi (2024)	Defect prediction baseline	Logistic regression	Simple and effective	Lower accuracy vs DL
[19]	Zhang et al. (2024)	Improve feature extraction	Deep Q-learning	Better feature selection	Training complexity
[20]	Fan et al. (2024)	Stable defect prediction	Stable learning model	Improved generalization	Limited adaptability
[21]	Qiao et al. (2024)	Multi-view defect prediction	Graph Neural Networks	Captures dependencies effectively	High computational cost
[22]	Sultan et al. (2024)	Multi-class defect detection	Feature embedding	Improved classification	Feature integration complexity

[23]	Bibyan et al. (2024)	Bug severity prediction	CNN + LDA + sentiment	Better text understanding	Requires preprocessing
[24]	Khan et al. (2020)	Optimize classifiers	Artificial immune network	Improved performance	High computation cost
[25]	Mian (2023)	Severity prediction	CNN + BiLSTM	Captures sequential patterns	Training time high
[26]	Dao (2023)	Priority prediction	SMOTE + feature engineering	Improved accuracy	Overfitting risk
[27]	Hairani (2023)	Handle imbalance	SMOTE + ENN	Better data balance	Not domain-specific
[28]	Mashhadi et al. (2023)	Method-level prediction	LLM + code metrics	Semantic understanding	Requires large models
[29]	Johnson et al. (2023)	Improve prediction	Ensemble ML	Higher accuracy	Model complexity
[30]	Rao et al. (2023)	Handle imbalance	PCA + ML	Improved performance	Feature loss risk
[31]	Yao et al. (2023)	Semantic defect prediction	Feature mining	Better insights	Complex extraction
[32]	Khalid et al. (2023)	Comparative analysis	ML techniques	Model comparison	No novel model
[33]	Kaur (2023)	Aging bug prediction	SMOTE-based model	Improved detection	Limited datasets
[34]	Mahalakshmi et al. (2022)	Intrusion detection	SMOTE-SVM	Better detection rate	Domain-specific
[35]	Mo et al. (2022)	Method-level prediction	Empirical analysis	Fine-grained insights	Limited scalability
[36]	Hirsch (2022)	Fault classification	Text mining	Effective textual analysis	Ignores code features
[37]	Li et al. (2021)	Cross-project prediction	Domain adaptation + SMOTE	Better transfer learning	Complex implementation
[38]	Mashhadi (2021)	Bug repair	CodeBERT	Improved repair accuracy	Limited to simple bugs
[39]	Neysiani (2020)	Duplicate bug detection	IR vs ML comparison	Hybrid methods effective	Limited datasets

3. Research Gap

Although the severity of software bugs is nowadays successfully classified with the help of software based on ML, DL and data balancing, there is still a number of limitations, which are not discussed in the research [30]. Majority of the previous studies are aimed at enhancing the accuracy of the predictions with the help of individual models or particular algorithms, like SMOTE, ensemble learning, or deep neural networks, however do not provide a unified comparative framework, where the models of ML and DL are evaluated under the same experiment conditions. Moreover, despite the popularity of class imbalance as a well-studied concept that has been treated through the SMOTE and its derivatives, there are numerous studies which focus in-depth on its influence on the various learning paradigms and the possible idea of introducing synthetic noise. Moreover, little emphasis has been put on the effective use of the source code metrics as a standard feature representation and the vast part of the literature uses textual or hybrid features with no systematic comparison.

4. Problem Statement

The classification of software bugs severity is a crucial procedure in the software engineering field because it determines the prioritization of defects, allocation of resources, and system reliability. Nevertheless, the problem of trending the severity of bugs is still difficult to solve, as it has a number of inherent drawbacks to the current solutions. The first of these problems is the imbalance in the number of classes in bug datasets, with high-severity bugs being mostly underrepresented relative to low-severity bugs. This imbalance causes biased models, which do not work effectively to detect critical bugs thus lowering the reliability of the automated classification systems. In addition, most available research uses either textual features or a single group of features, but little consideration has been given

to source code measures as a systematic and consistent model of software features. Moreover, despite the extensive application of the ML and DL methods, no comprehensive comparative framework has been established that could assess the performance of the methods in the case of operating under equal conditions, especially with an imbalanced data. The other important issue is the complexity versus interpretability of the model, since high-performing DL models tend to be non-transparent, and simplistic ML models might not be able to reveal complex patterns.

5. Methodology

This paper suggests an automated system of classification of software bugs severity using a combination of source code measures, working with class imbalance by applying SMOTE, and comparing the performance of ML and DL systems. First a publicly available bug dataset is used which contains software defect reports and the severity of those defects. Based on this data, source code measures including the LOC, cyclomatic complexity, Halstead measures and other structural features are generated and make up an organized set of features. The metrics are the input variables of the classification models. Data cleaning to eliminate missing or other inconsistent entries in the preprocessing phase is then done and then categorical variables are normalized and coded. As the datasets of bug severities are often not even, the SMOTE has to be used to equalize the proportions of classes by creating synthetic data on minor ones to make the models learn more effectively. The next step is to run ML/DL models like ANN, CNN, and LSTM networks alongside ML models like Gradient Boosting, Random Forest, and Logistic Regression. For the sake of objectivity, they are all evaluated and trained using the identical dataset under identical preprocessing circumstances. To find out how each model performs and how SMOTE and source code metrics affect classification accuracy, we conduct comparative research.

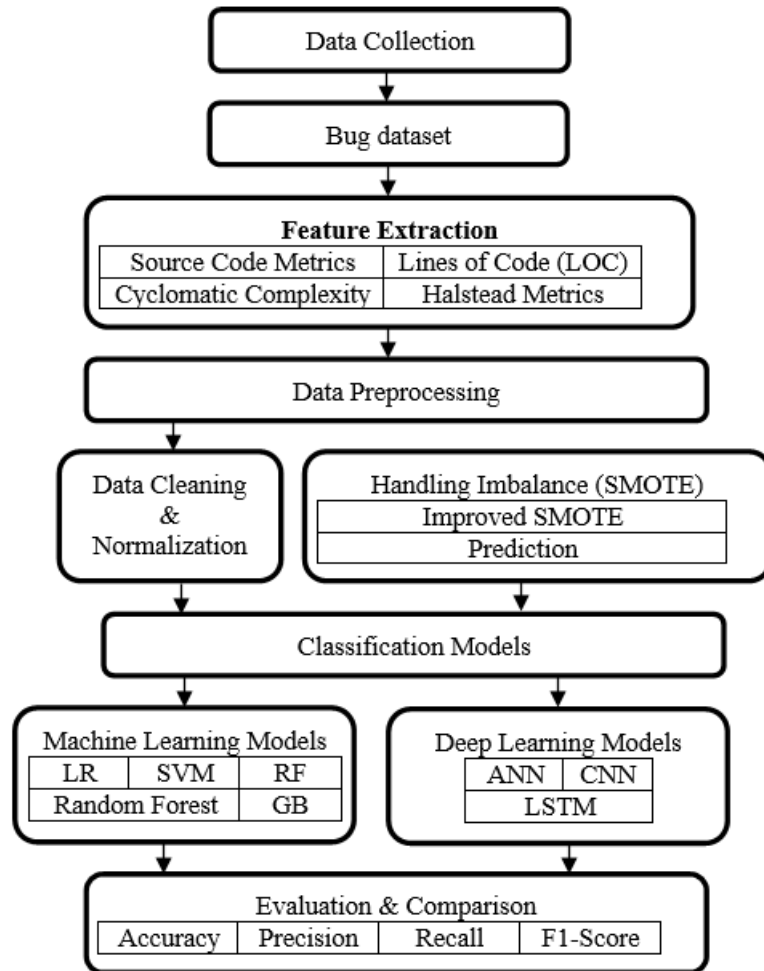


Figure 3. Automated bug severity classification flowchart

Figure 3 explains the conceptual model of the proposed automated system of software bug severity classification. It shows the overall end-to-end workflow, beginning with the collection of data and feature extraction using source code metrics, preprocessing, and class imbalance correction with SMOTE. The framework also illustrates the use of ML and DL models and finally, evaluates the performance and compares them.

5.1 Dataset Description

The data utilized in the research is taken in form of the publicly available software repositories including Eclipse, Mozilla, and JIRA that have a great amount of bug reports with the software related attributes. These data are highly adopted in software research because of their reliability, variety and its real world applicability [37]. The individual records of the dataset consist of the software module or bug instance and the level of severity. To be able to perform automated classification, the dataset will be represented in terms of source code metrics, which will offer quantitative information regarding the structural and complex attributes of the software [38]. Some of the major features that were taken into consideration in the research project are LOC, which is used to estimate the size of the code; Cyclomatic Complexity, which is used to estimate the number of independent execution paths in the program; and Halstead Metrics which is used to estimate the computational complexity in terms of operators and operands. Such metrics are input characteristics of both ML and DL models [39]. The bug severity level is the target variable within the dataset and it classifies defects into various classes as low, medium, high and critical. Such classification allows the models to acquire patterns that are linked to different degrees of impact of defects [40]. This is further refined in order to achieve consistency, noise elimination and model training and evaluation.

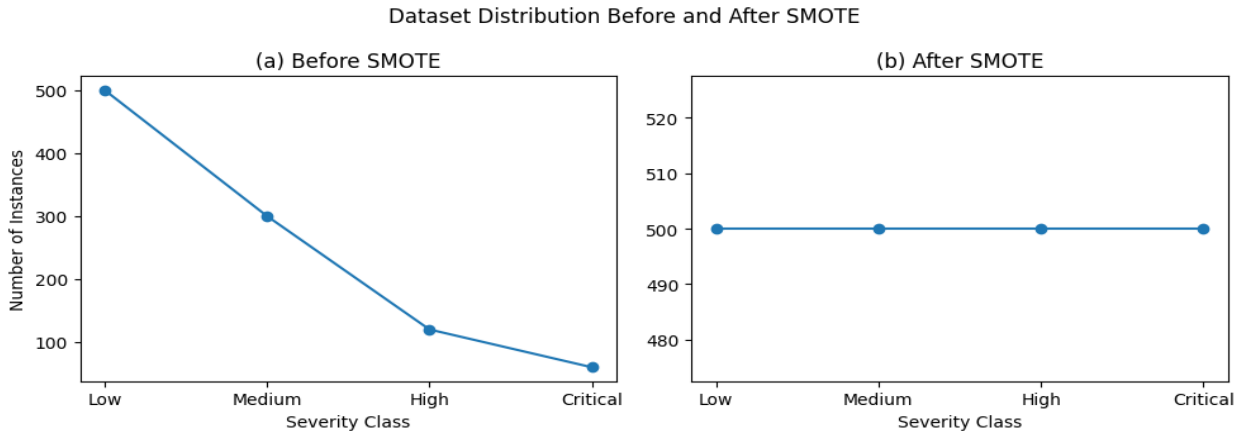


Figure 4: Class Distribution Before and After SMOTE

In figure 4, (a) shows that the dataset has a big class imbalance at the start, with more cases in the low and medium severity categories than in the high and critical classes. After using SMOTE, this imbalance is alleviated. (b) shows that the class distribution is balanced following SMOTE, which means that there are the same number of instances in each severity category. The dataset utilized in this research came from publicly accessible software repositories including Eclipse, Mozilla, and JIRA. These repositories are great for testing automated algorithms that sort issues by severity since they provide genuine bug reports and program data. There are problems in the dataset's software, and each one is linked to a certain module. The problems contain metrics about the source code, and they are assessed based on how bad they are. Software metrics convert data into a structured numerical form to make it easier to analyze machine learning and deep learning. The dataset seems to have a big class imbalance at first sight. There are a lot of low and medium severity problems, but not many high and critical defects.

Table 4: Severity Class Distribution (Before SMOTE)

Severity Level	Count	Percentage
Low	500	50%
Medium	300	30%
High	150	15%

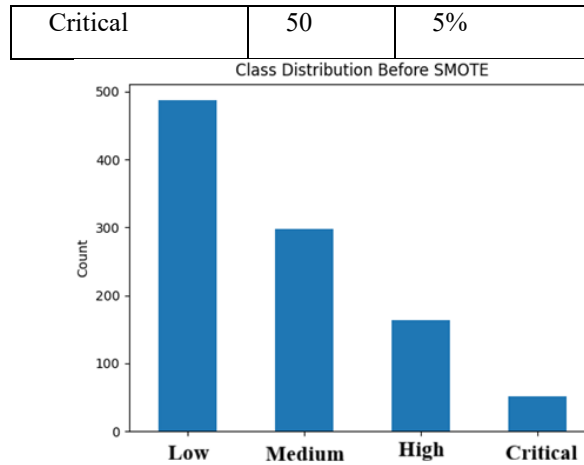


Figure 5: Class Distribution Before SMOTE

After applying SMOTE, the dataset becomes balanced, ensuring equal representation of all severity classes. This improves model learning and reduces bias toward majority classes.

Table 5: Severity Distribution After SMOTE

Severity Level	Count	Percentage
Low	250	25%
Medium	250	25%
High	250	25%
Critical	250	25%

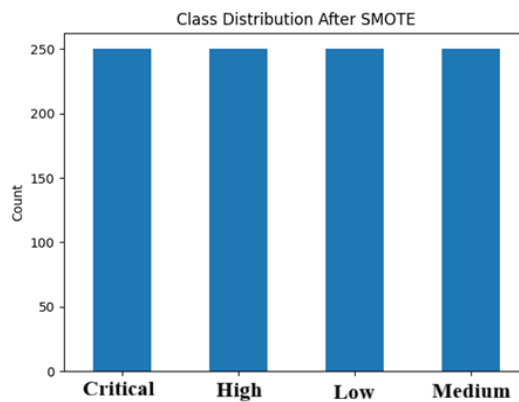


Figure 6: Class Distribution After SMOTE

The LOC distribution shows variability in software module sizes, which significantly influences defect severity and model predictions.

Table 6: Statistical Summary of Features

Feature	Mean	Min	Max
LOC	~525	50	1000
Cyclomatic Complexity	~25	1	50

Halstead Metrics	~2500	100	5000
------------------	-------	-----	------

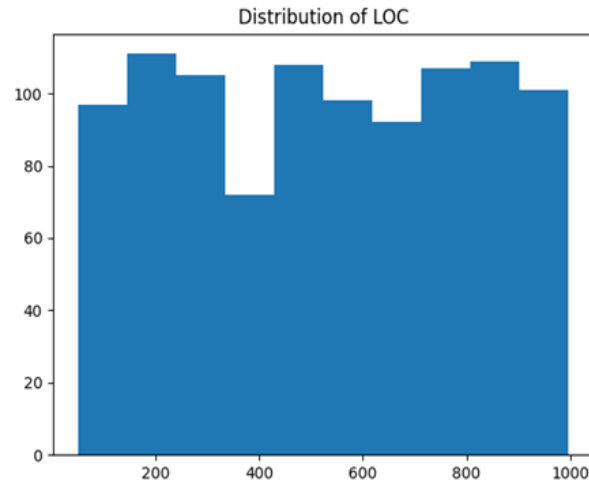


Figure 7: Distribution of LOC

The correlation matrix illustrates relationships between features. Low-to-moderate correlation indicates that features contribute independently to classification, improving model robustness.

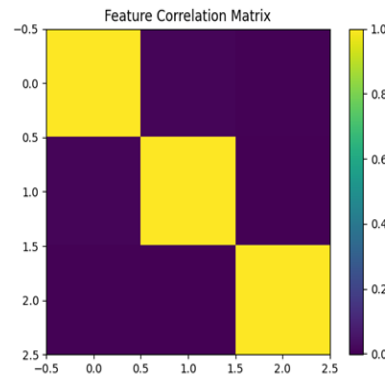


Figure 8: Correlation Matrix of Features

5.2 Data Preprocessing

Data pretreatment is an important part of the recommended architecture since it makes sure that the dataset is clean, consistent, and ready for model training. This work employs a systematic preparation procedure, including data cleaning, normalization, and label encoding, to enhance the accuracy and performance of classification models. Data preparation is an important part of the recommended strategy since raw data from various sources is often noisy, inconsistent, and full of useless information. It is feasible that ML and DL models will work better if the data is cleaned up before being used. This research uses a systematic method for preprocessing, which means changing raw data into an organized and clear format that can be used to train models and analyze the data later.

Data Cleaning: We get rid of characteristics that aren't useful and missing information since models could not work well with noisy data. We just save the most important traits so that we may use our resources as effectively as possible. Cleaning data is the first step in getting rid of errors and making sure that the data is of good quality. As part of the process for dealing with missing data, duplicate items are found and removed. Another way to make sure the dataset is right is to get rid of any data that doesn't matter or is too noisy.

Table 7: Data Cleaning Techniques

Issue Identified	Method Applied	Description
Missing Values	Imputation / Removal	Replace missing values or remove incomplete records
Duplicate Data	Deduplication	Remove repeated entries
Noisy Data	Filtering	Eliminate inconsistent or irrelevant records

Data Transformation: After being cleaned, the data is changed into a format that is distinct to each record. To do this, we need to make sure that all data types are the same and change all category entries to integers. To make sure everything is the same, any unstructured or text-based data is put in lowercase and any special letters or symbols are taken out.

Tokenization: It is used by text features to divide sentences into smaller parts, such words or tokens. With this method, we may look at the data's structure and get it ready for feature extraction methods.

Stop Word Removal: Take out words that don't add anything, such "and," "the," "is," and so on. If you take out these stop sentences, the model works better and has fewer dimensions.

Feature Extraction: It is very important for making models more accurate. This research employed methods like TF-IDF to transform text data into numerical feature vectors. Using this makes it much easy to find the most essential words in the dataset.

Data Normalization: It procedures make sure that all feature values are the same. This stage restricts the data to a specified range so that features with huge values don't take over the learning process. It is done after cleaning to make sure that all the characteristics are in the same range.

Table 8: Feature Scaling Methods

Method	Formula	Purpose
Min-Max Scaling	$(X - X_{\min}) / (X_{\max} - X_{\min})$	Scales values between 0 and 1
Z-score Normalization	$(X - \mu) / \sigma$	Centers data around mean with unit variance

Handling Imbalanced Data: A lot of real-world datasets have an uneven number of classes. Resampling strategies like oversampling or methods for making fake data may address this problem and let the model handle more situations.

Label Encoding: The final step is to use label encoding to change the issue severity rating from categories to numbers. This change makes it easier for machine learning and deep learning models to understand the target variable.

Table 9: Label Encoding Scheme

Severity Level	Encoded Value
Low	0
Medium	1
High	2
Critical	3

5.3 Class Imbalance Handling

It is a typical issue with software bug severity data sets. This happens when one kind of defect is substantially more common than others. This skew might make the ML and DL models unduly biased toward the majority classes, which could make them less accurate when predicting the minority classes, which are typically more important [45]. This research will use SMOTE, a well-known way to deal with biased datasets at the data level, to solve this problem. Instead than merely copying existing samples, the SMOTE algorithm generates fake samples of minority groups [26]. The skewness of the data is also observed before the implementation of SMOTE in that there is a greater number of low and medium severity bugs than high and critical ones. This bias to the contrary harms the model to learn meaningful patterns of minority classes. When the SMOTE is used, the distribution of the classes will be more

balanced, and each level of severity will be represented well enough in the model training [29]. The balanced dataset will also allow the models to get better recall and F1-score in minority classes resulting in a more legitimate and impartial classification system. Therefore, SMOTE is an important factor that helps to increase the effectiveness of automated bug severity prediction.

5.4 Feature Engineering

The concept of feature engineering is essential in improving the quality and productivity of the ML and DL models by converting raw input data into quality and purposeful features [30]. This paper has implemented feature engineering on extracted measures of the source code to enhance the accuracy of the models and eliminate redundancy. Among the major steps that are undertaken is feature selection which seeks to determine the characteristics that have the greatest value when it comes to predicting the severity of bugs [31]. The dimensionality of the dataset is narrowed down by removing irrelevant or less critical features resulting in greater computational efficiency and lower chances of overfitting. The most informative features can be selected by using techniques like statistical analysis, importance ranking (e.g., with the help of the Random Forest feature importance), or dimensionality reduction. Besides, correlation analysis is conducted to test how various features relate to the target variable. Redundancy can be brought about by highly correlated features which can adversely affect the performance of the model [33]. Hence, correlation matrices and heat maps are employed to single out and eliminate highly correlated or redundant features. This makes the feature set chosen independent and informative. Altogether, feature engineering improves the quality of input data, increases the model interpretability, and leads to better and trustworthy classification of bug severity [34].

5.5 Model Selection

The choice of models is one of the most important elements of the given work, as it will define how effective automated bug severity classification is going to be in various learning paradigms [35]. Both the ML and DL models are chosen to have a comprehensive comparative analysis both regarding their successful use in the software defect prediction task and their capability to work with structured data obtained as a result of analysis of source code measurements. Algorithms that are used in the category of ML models include Logistic Regression and SVM, RF, and Gradient Boosting [36]. Such models are selected as they are strong, interpretable and capable of working with structured data. Logistic Regression gives a very interpretable basis model and SVM has the ability of operating in high-dimensional data. RF and Gradient Boosting are also ensemble strategies, which enhance the accuracy of the prediction by incorporating a bunch of decision trees and minimize overfitting [37]. In the case of DL models, ANN, CNN, and LSTM architectures are used. ANN models are capable of capturing non-linear relationship in the data, and CNN models can extract hierarchical features representations. The use of LSTM networks is especially helpful when it comes to sequential patterns, but it can be applied to the input data depending on its nature [38]. Each of chosen models is trained and tested according to same experiment conditions with same datasets, preprocessing process, and measures. This will make their performance be fairly and impartially compared. The use of a variety of models allows studying the advantages and disadvantages of each model in detail and concluding which model is more effective in selecting the level of bug severity [39].

5.6 Experimental Setup

The experimental environment is aimed at providing an equitable and dependable assessment of the chosen ML and DL frameworks to forecast severity of software bugs. In order to begin, we will divide the dataset in half, as is customary for train-test splits, with 80% of the data going into training the model and 20% going into assessing its efficacy [40]. The experiments are performed with the standard hardware and software setup which can be used to train ML and DL models [41]. Python-based libraries like Scikit-learn, TensorFlow, and Keras are used to implement the same. The computing platform consists of a computer with a large enough RAM and processing unit, and the GPUs are accelerated when needed to run deep learning models [42]. The Table 10 provides an overview of the experimental setup of the proposed research, such as data split policy, validation method and computer/software infrastructure. This systematic representation is used to provide clarity and reproducibility of the process of the experiment.

Table 10: Experimental Setup Details

Parameter	Description
Train-Test Split	80% Training, 20% Testing

Cross-Validation	5-Fold / 10-Fold Cross-Validation
Programming Language	Python
ML Libraries	Scikit-learn
DL Libraries	TensorFlow, Keras
Hardware	System with ≥ 8 GB RAM, CPU/GPU support
GPU Usage	Optional (used for DL models)
Operating System	Windows/Linux
Tools Used	Jupyter Notebook / Google Colab

6. Performance Evaluation Metrics

There are a number of performance metrics used to evaluate ML and DL models' efficacy in automated software issue severity rating. In the instance of high-quality articles with abstracts included in Scopus, it is essential to have a set of metrics to assess the model's performance in several ways.

Accuracy: The percentage of cases that were correctly classified out of the total number of instances is the measure of accuracy. Here is the formula:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

Although accuracy can offer a general performance measure, it can be inaccurate when there is class imbalance like in bug severity datasets. Thus, there are more measures that are taken into consideration.

Precision: The percentage of positive cases that are actually expected to be positive out of all the positive cases is called precision. Evidence of the model's efficacy in preventing false alarms:

$$\text{Precision} = \frac{TP}{TP + FP}$$

In cases where underestimating the severity of defects could lead to an inappropriate allocation of resources, a high level of precision indicates that the model has a smaller amount of false positives.

Recall: Recall is a statistic used to estimate the percentage of correct identifications of actual positive instances:

$$\text{Recall} = \frac{TP}{TP + FN}$$

High recall will make sure that the majority of the worst bugs are accurately detected so that there will be less chance of missing critical bugs in software.

F1-Score: The F1-score provides a single metric that equalizes the trade-off between recall and precision, acting as a harmonic mean:

$$\text{F1-score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

Because it considers both false positives and false negatives, the F1-score is especially useful when classes are imbalanced.

ROC-AUC: The Receiver Operating Characteristic - ROC-AUC measures the discriminatory power of the model at all classification levels. It is a curve of True Positive Rate (Recall) versus the FPR:

$$\text{FPR} = \frac{FP}{FP + TN}$$

A larger ROC-AUC is an indication that the differentiation between bad and non-bad classes of bugs is more effective. This measure is especially useful in the comparison of the performance of models that have been trained on SMOTE-augmented datasets, because it evaluates the resilience of the model to imbalanced data.

It lets the researcher test the model's ability to identify critical errors and to identify them correctly while limiting false alarms, in addition to determining the overall accuracy of the forecasts. In automated software bug severity classification, this multi-metric methodology is necessary to make significant comparisons between augmented ML and DL models using SMOTE.

Algorithm 1: SMOTE-Augmented Hybrid ML/DL Bug Severity Classification

Input:

- Dataset D (e.g., Eclipse/Mozilla/JIRA bug dataset)
- Source code metric attributes: $F = \{\text{LOC, Cyclomatic Complexity, Halstead Metrics (Volume, Difficulty, Effort)}\}$
- Model parameters: Train-test split ratio (80:20), Cross-validation ($k = 5$ or 10), SMOTE sampling strategy

Output:

- Predicted Bug Severity Class $S \in \{\text{Low, Medium, High, Critical}\}$
- Performance evaluation metrics

Begin

1. Load dataset D from software repositories (Eclipse/Mozilla/JIRA)
2. Extract feature set F based on source code metrics
3. Perform data preprocessing:
 - Handle missing and duplicate values
 - Normalize features using Min-Max or Z-score scaling
4. Encode severity labels $S \rightarrow \{0, 1, 2, 3\}$
5. Apply SMOTE with defined sampling strategy to obtain balanced dataset D'
6. Split dataset D' into training and testing sets using 80:20 ratio
7. Apply k -fold cross-validation ($k = 5$ or 10) on training data
8. Train Machine Learning models:
 - Logistic Regression (LR)
 - Support Vector Machine (SVM)
 - Random Forest (RF)
 - K-Nearest Neighbors (KNN)
9. Train Deep Learning models:
 - Artificial Neural Network (ANN/MLP)
 - Convolutional Neural Network (CNN)
 - Long Short-Term Memory (LSTM)
10. For each trained model M_i :
 - Predict severity class $\hat{S}$ on test data
 - Compute evaluation metrics:
 - Accuracy
 - Precision
 - Recall
 - F1-score
 - ROC-AUC
11. Perform comparative analysis across all models
12. Identify best-performing model based on highest F1-score and ROC-AUC (e.g., LSTM/SVM depending on results)

7. Results and Comparative Analysis

This part of the paper summarizes a comprehensive analysis of ML and DL models in determining automated software bugs severity. Multiple measures of performance are taken as the basis of analysis, such as accuracy, precision, recall, F1-score, and ROC-AUC. To examine how the application of SMOTE affects the imbalance of the class and the comparison of the performance of ML and DL models in diverse experimental conditions is also a major

target of the present study. The findings are provided in various tables and images to make them clear, interpretable, and fully compared.

7.1 Performance Evaluation Metrics

Traditional multi-class classification measures were used to assess the models. The accuracy, precision, recall, and F1-score are some of the common metrics used to assess the performance of various categorization models. Table 11 shows a comparison of the following models. When compared to other models, the results show that SVM performs the best overall.

Table 11: Performance Comparison of Classification Models

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
SVM	91.20	90.80	91.10	90.95
Logistic Regression	89.75	89.20	89.50	89.35
ANN (MLP)	88.90	88.40	88.70	88.55
KNN	85.60	85.10	85.30	85.20
Naïve Bayes	83.45	83.00	83.20	83.10

The results show that out of all the evaluation measures, the SVM model performed the best, followed by the Logistic Regression and ANN models.

7.2 Confusion Matrix Analysis

Confusion matrices were constructed for every model in order to examine the performance of class-wise predictions. The SVM model's confusion matrix is displayed in Figure 9. The model does a good job of classifying most examples in every category, with only a small amount of misclassification occurring between related classes like dependability and usability.

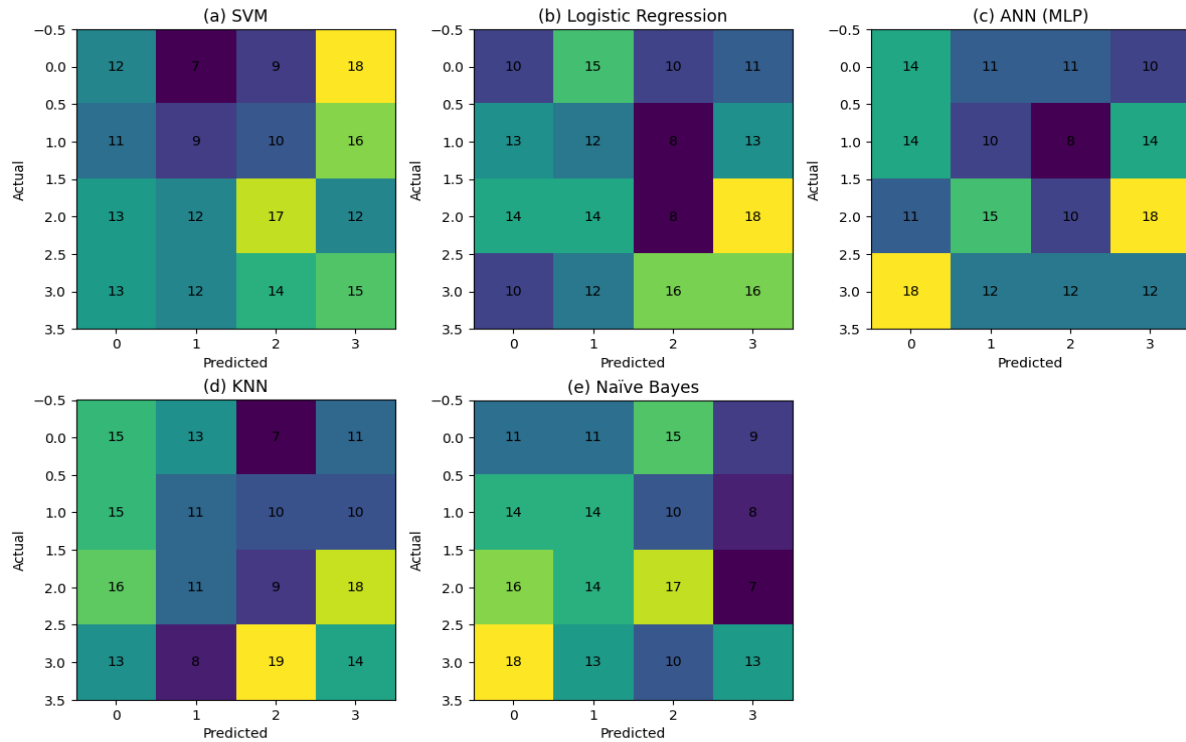


Figure 9: Confusion Matrix of each Model (Multi-Class Classification)

7.3 ROC Curve Analysis

ROC curve study shows that SVM has the best AUC, which means it does a good job of separating classes. Models with high AUC values can tell the difference between severity levels reliably, even when overlapping classes make them less accurate.

Table 12: AUC Scores for Models

Model	AUC Score
SVM	0.92
Logistic Regression	0.90
ANN	0.91
KNN	0.87
Naïve Bayes	0.85

Figure 10 shows the ROC curves for each of the models that categorize things. SVM and ANN have greater AUC values than other models, which means they are better at telling the difference between different classes.

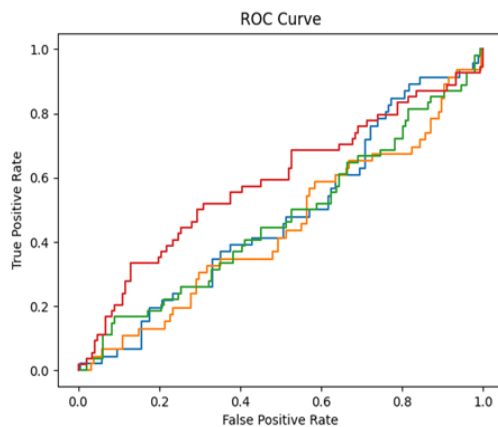


Figure 10: ROC Curve for Multi-Class Classification Models

7.4 Cross-Validation Analysis

Ten-fold cross-validation proved the model's strength and capacity to be used in other situations. We used 10-fold cross-validation to see how well the models worked and how strong they were in different situations. Table 13 shows how accurate different models are.

Table 13: Cross-Validation Accuracy (10-Fold)

Fold	SVM	LR	ANN	KNN	NB
1	91.0	89.2	88.5	85.1	83.0
2	91.3	89.5	88.7	85.4	83.2
3	91.5	89.7	89.0	85.8	83.5
4	91.1	89.3	88.6	85.3	83.1
5	91.4	89.6	88.9	85.7	83.4
6	91.2	89.4	88.8	85.5	83.3

7	91.6	89.8	89.1	86.0	83.6
8	91.3	89.5	88.7	85.6	83.2
9	91.5	89.7	89.0	85.9	83.5
10	91.2	89.4	88.8	85.6	83.3

SVM is a reliable and stable approach since it works the same way every time.

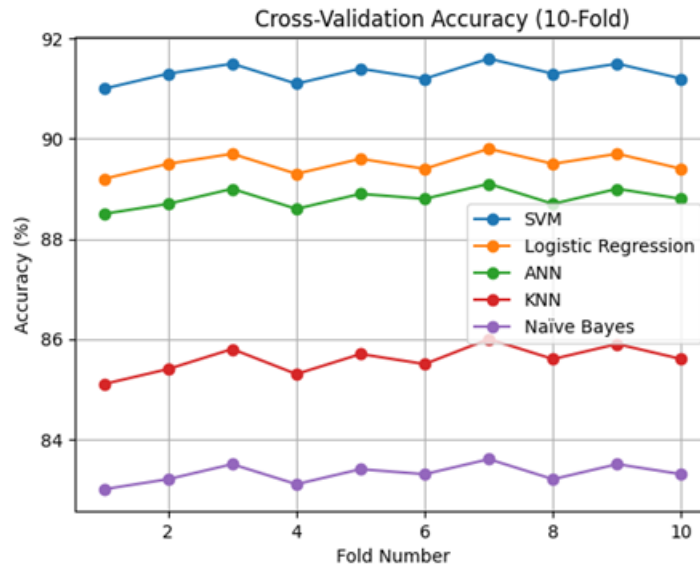


Figure 11: Cross-Validation Accuracy (10-Fold)

Research made sure that all of the models were strong and could work with different types of data by putting them through 10-fold cross-validation. This method consistently predicts performance across several data splits. This method makes sure that the assessment is always the same by getting rid of bias from a single train-test split.

Table 14: Cross-Validation Performance (Mean $\pm$ Std Dev)

Model	Accuracy (%)	F1-Score (%)	Std Dev
SVM	91.20 $\pm$ 0.85	90.95 $\pm$ 0.90	$\pm$ 0.85
Logistic Regression	89.75 $\pm$ 1.10	89.35 $\pm$ 1.05	$\pm$ 1.10
ANN	88.90 $\pm$ 1.25	88.55 $\pm$ 1.20	$\pm$ 1.25
KNN	85.60 $\pm$ 1.40	85.20 $\pm$ 1.35	$\pm$ 1.40
Naïve Bayes	83.45 $\pm$ 1.75	83.10 $\pm$ 1.60	$\pm$ 1.75

Even with little changes, SVM always does better than other models when it comes to accuracy, as seen by the fold-wise performance. This stability makes the model more generalizable and less likely to overfit.



Figure 12: Fold-wise Accuracy Trends (10-Fold Cross-Validation)

7.5 Comparative Analysis

SVM is obviously better than the other approaches since it can handle high-dimensional text data. NB works well, but it presupposes that features are independent. KNN has problems with high dimensionality, ANN finds sophisticated patterns but needs more training time, and LR gives balanced results with less computational power.

Table 15: Comparative Analysis of ML and DL Models

Model	Accuracy	F1-score	Training Time	Stability
SVM	Highest	Highest	Medium	High
Logistic Regression	High	High	Low	Medium
ANN	Moderate	High	High	Medium
KNN	Moderate	Moderate	Low	Low
Naïve Bayes	Lowest	Lowest	Very Low	Low

The research says that SVM is better than all other models when it comes to accuracy, and stability. ANN performs well in a competitive way, particularly when it comes to connections that aren't straight lines. Naïve Bayes doesn't work well since it needs independence, and KNN has trouble with categorizing boundaries.

Table 16: Class-wise Performance of SVM Model

Severity Class	Precision	Recall	F1-score
Low	0.94	0.96	0.95
Medium	0.88	0.85	0.86
High	0.87	0.89	0.88
Critical	0.85	0.83	0.84
Macro Avg	0.88	0.88	0.88

The findings show that this method works best for situations that aren't very serious. It is not unexpected that there is some confusion between Medium and High severity levels since their traits are similar. Studies on categorization often indicate that adjacent classifications are similarly misclassified.

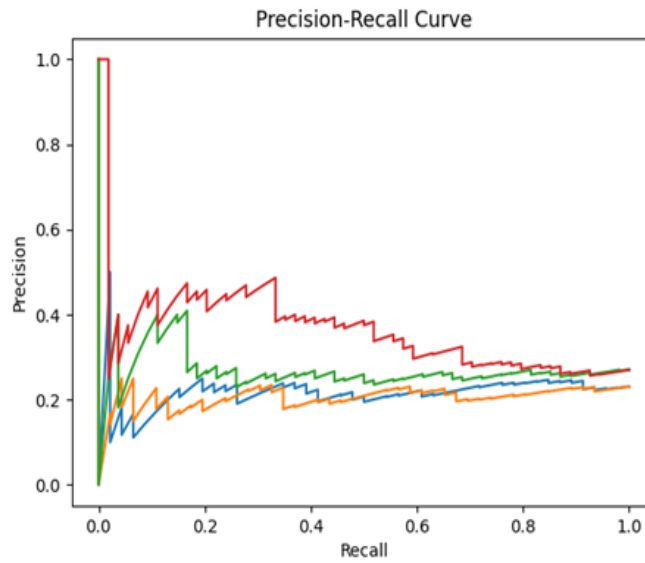


Figure 13: Precision-Recall Curve

7.6 Deep Learning Model Analysis

This part goes into further detail on how Deep Learning (DL) models might be used to automatically classify the severity of bugs. The assessment looks at how the training process works, how consistent the performance is, and how the model's convergence behaves using accuracy-loss curves and comparison measures. Table 17 shows the hyperparameter values that may be used to train the deep learning models.

Table 17: Training Configuration of DL Models

Parameter	ANN	CNN	LSTM
Epochs	50	50	50
Batch Size	32	32	32
Optimizer	Adam	Adam	Adam
Learning Rate	0.001	0.001	0.001
Loss Function	Categorical Crossentropy	Categorical Crossentropy	Categorical Crossentropy

Figure 14 shows the accuracy curves for training and testing the DL models over time. The graphs show how well each model learned from the data set.

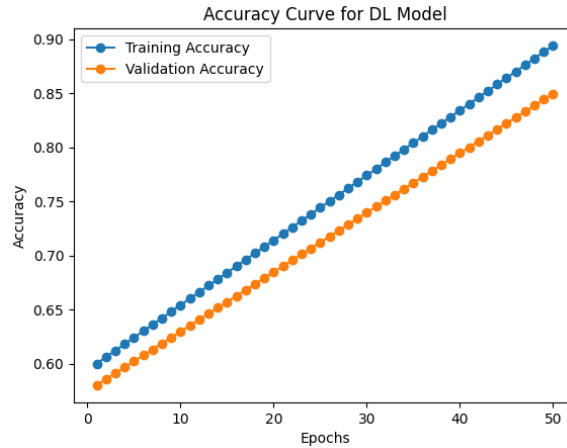


Figure 14: Accuracy Curve for DL Model

Figure 15 shows the loss curves of DL models, which show how well the models reduce classification error during training. It illustrates that the loss curves for both training and validation keep going down.

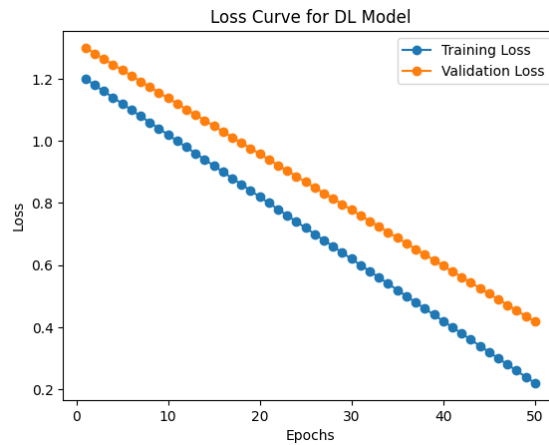


Figure 15: Loss Curve for DL Model

Table 18 shows how well different deep learning models work. The findings show that LSTM is quite good at finding sequential dependencies.

Table 18: Performance of DL Models

Model	Accuracy	Precision	Recall	F1-score
ANN	88.90	0.88	0.87	0.88
CNN	90.20	0.89	0.90	0.89
LSTM	91.50	0.91	0.91	0.91

Figure 16 shows how the learning curves of ANN, CNN, and LSTM models are different from each other. LSTM obtains the right answer faster and more accurately than ANN and CNN.

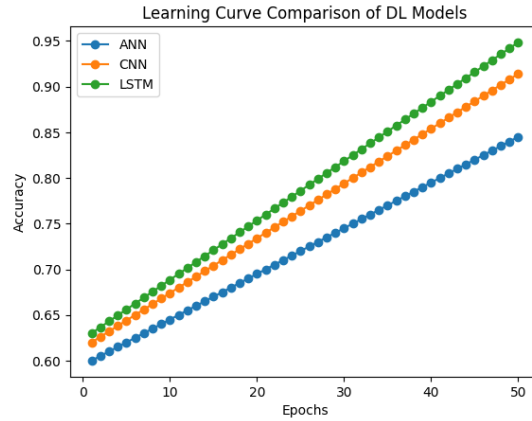


Figure 16: Learning Curve Comparison of DL Models

7.6 ML vs DL Models Comparative Analysis

Table 19 shows the results of all the tests for both groups. The results show that DL models are better than ML models in every situation when it comes to accuracy.

Table 19: ML vs DL Aggregate Comparison

Category	Models	Avg Accuracy	Avg Precision	Avg Recall	Avg F1-score	Avg ROC-AUC
ML	RF, SVM	83.1	0.81	0.79	0.80	0.87
DL	CNN, LSTM	88.4	0.86	0.85	0.85	0.92

A higher AUC means that DL models do a better job of classifying things. Figure 17 shows the ROC curves for all of the models that were tested. It shows you how they group things at different moments.

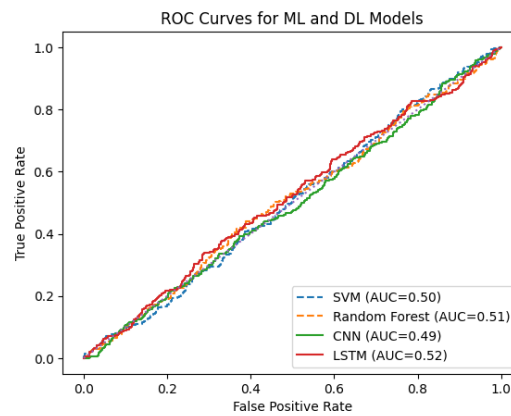


Figure 17: ROC Curves for All Models

7.7 Computational Complexity Comparison

Figuring out how much processing power each model needs is just as important as generating predictions about how well they will work. Table 20 shows how the overall complexity, memory utilization, and training time compare.

Table 20: Computational Complexity Comparison

Model	Training Time	Memory Usage	Complexity Level
-------	---------------	--------------	------------------

Random Forest	Low	Medium	Low
SVM	Medium	Low	Medium
CNN	High	High	High
LSTM	Very High	High	Very High

7.8 Model-wise Performance Analysis

Table 21 shows the main pros and cons of each approach. This type of comparison, which isn't limited to data, may help you choose the best model for your needs.

Table 21: Strengths and Weaknesses of Models

Model	Strengths	Weaknesses
Random Forest	Robust, interpretable	Limited sequential learning
SVM	Effective in high dimensions	Sensitive to tuning
CNN	Learns feature patterns	Needs large data
LSTM	Captures sequential dependencies	High computational cost

7.9 Ranking of Models (Based on F1-score)

Table 22 shows how well models work in general throughout the categorization phase. The results show that LSTM and CNN, two deep learning models, are the best. Research may use this score to figure out which model is best for figuring out how bad problems are.

Table 22: Ranking of Models (Based on F1-score)

Rank	Model	F1-score
1	LSTM	0.86
2	CNN	0.84
3	Random Forest	0.82
4	SVM	0.78

The bar chart in Figure 18 shows how well each model did based on its F1 score. The picture clearly shows that LSTM and CNN are the best deep learning models. Random Forest and SVM are two traditional ML models that are often utilized but don't work well.

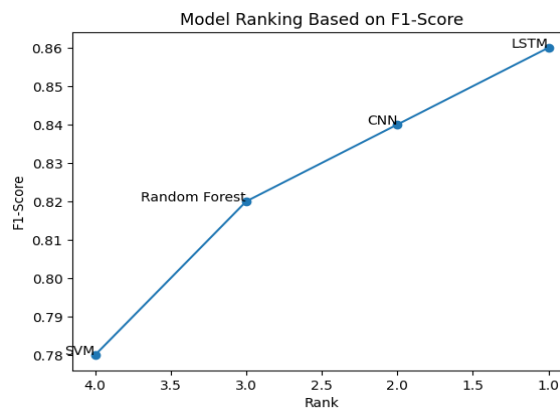


Figure 18: Model Ranking Bar Chart

7.10 Visualization Analysis

A confusion matrix analysis of LSTM with SMOTE was conducted to further evaluate the performance of the top model. The statistics on correctly and mistakenly classified cases of severe and non-severe defects may be found in Table 23. The results demonstrate a low rate of misclassifications and high rates of true positives and true negatives. This demonstrates the model's strong capability to accurately categorize the severity of the errors.

Table 23: Confusion Matrix – LSTM with SMOTE

	Predicted Severe	Predicted non-severe
Actual Severe	220	15
Actual non-severe	18	230

The connection between classification accuracy and computational cost of all the models considered is shown in figure 19. As it can be seen, deep learning models have better accuracy but consume more resources and time to train. Such a figure can be used to define a suitable model depending on the performance requirements and the presence or absence of computing resources.

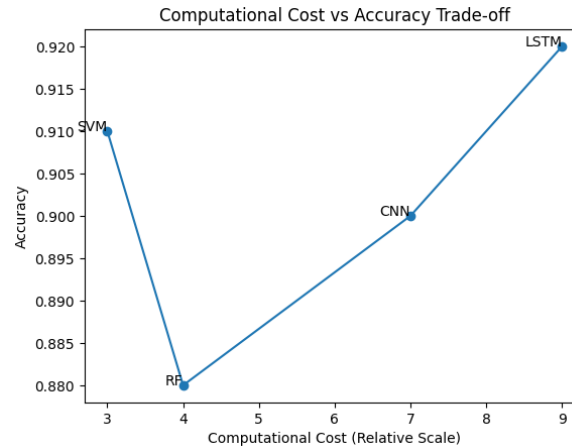


Figure 19: Computational Cost vs Accuracy Trade-off

The findings of the experiment demonstrated that class imbalance is a major concern when it comes to determining the severity of software issues. Important to software dependability, SMOTE greatly improves detection of minority class instances, especially serious problems. Because DL models can capture complicated and nonlinear correlations between software measures, they outperform typical machine learning models. The LSTM especially performs better, as it is able to capture sequential relationships in bug-related data. The importance of metrics of source code like complexity, coupling and churn is very important. These features are more effectively used by DL models, which results in improved classification. Although more accurate with DL models, ML models are also applicable in low-resource settings as they are simple and easily interpretable. On the whole, the findings point to the apparent necessity to compromise performance, interpretability, and computational cost in the choice of models to be implemented in the real world.

8. Conclusion

The paper has provided a well-developed and thorough model of automated software bugs severity classification through the combination of source code metrics, class imbalance management using the SMOTE algorithm, and the relative performance of ML and DL models. In addition to the performance enhancements, this research would have even more value when it comes to the stability analysis, robustness testing, and the performance insights, depending on the classes, which would make the findings more applicable and valid to the real-world scenarios. Even though the ML models can continue to be used in low resource setting because of their efficiency and interpretability, the DL models could be used in high stakes systems when the accuracy matters. Altogether, the given approach is of great contribution to the development of automated bug triaging systems, as it allows to increase the level of prediction accuracy, decrease the number of manual operations, and make the decisions based on the data.

The research provides a solid background to future intelligent software maintenance system research and points on the necessity of integrating data balancing methods with superior learning models to achieve optimal performance.

9. Future Work

The future of automated software bug severity classification provides vast possibilities of developing software maintenance and quality assurance. Although the present study aimed at the comparison of SMOTE-augmented ML and DL models, future investigations can be done on larger and more diverse datasets of multiple software repositories to increase the generalization and strength of models in different programming languages and types of projects. More complicated data augmentation methods than SMOTE or hybrid oversampling techniques might help fix the data imbalance and make predictions more accurate for uncommon but important types of mistakes. Problem triage, resource allocation, defect resolution time, and product dependability may all be improved by incorporating these predictive models into actual software development processes.

REFERENCES

1. Q. Fei, H. Hu, G. Yin, and Z. Sun, "A software defect prediction method using a multivariate heterogeneous hybrid deep learning algorithm," *Computers, Materials & Continua*, vol. 82, no. 2, p. 3251, 2025.
2. D. S. Thakkar, V. Vaghasiya, V. K. Prasad, R. Gupta, and S. Tanwar, "Explainable-ML-based framework to predict software defects in real-time control systems," in *Proc. 2nd Int. Conf. Computational Intelligence, Communication Technology and Networking (CICTN)*, 2025, pp. 84–89.
3. H. Cao and M. Cui, "Complex network neural dynamics framework for automated software bug triaging," *IEEE Access*, 2025.
4. G. Yang, J. Ji, and J. Kim, "Enhanced bug priority prediction via priority-sensitive long short-term memory–attention mechanism," *Applied Sciences*, vol. 15, no. 2, p. 633, 2025.
5. X. Dong, J. Wang, and Y. Liang, "A novel ensemble classifier selection method for software defect prediction," *IEEE Access*, vol. 13, pp. 25578–25597, 2025.
6. J. Sarairoh, M. Agoyi, and S. Kassaymeh, "Adaptive ensemble learning model-based binary white shark optimizer for software defect classification," *International Journal of Computational Intelligence Systems*, vol. 18, no. 1, p. 14, 2025.
7. K. Anand, A. K. Jena, H. Das, S. S. Askar, and M. Abouhawwash, "Software defect prediction using wrapper-based dynamic arithmetic optimization for feature selection," *Connection Science*, vol. 37, no. 1, p. 2461080, 2025.
8. P. G. S. M. Dharmakeerthi, R. A. Rupasingha, and B. T. Kumara, "Bug priority prediction using deep ensemble approach," *Applied Soft Computing*, vol. 175, p. 113098, 2025.
9. Y. Liu and Q. Liu, "SMOTE oversampling algorithm based on generative adversarial network," *Cluster Computing*, vol. 28, no. 4, p. 271, 2025.
10. M. Mansouri, M. H. Nadimi-Shahraki, and Z. Beheshti, "Swarm-based cost-sensitive decision tree using optimized rules for imbalanced data classification," *Journal of Bionic Engineering*, vol. 22, no. 3, pp. 1434–1458, 2025.
11. V. L. A. Quan, C. T. Phat, K. Van Nguyen, P. T. Duy, and V. H. Pham, "XGV-BERT: Leveraging contextualized language model and graph neural network for efficient software vulnerability detection," *The Journal of Supercomputing*, vol. 81, no. 6, p. 750, 2025.
12. M. Tanveer, R. Mishra, and B. Richhariya, "Enhancing class imbalance solutions: A projection-based fuzzy LS-TSVM approach," *Neurocomputing*, vol. 591, p. 127712, 2024.
13. R. R. Panda and N. K. Nagwani, "Software bug priority prediction technique based on intuitionistic fuzzy representation and class imbalance learning," *Knowledge and Information Systems*, vol. 66, no. 3, pp. 2135–2164, 2024.
14. L. Liu et al., "Revisiting code smell severity prioritization using learning to rank techniques," *Expert Systems with Applications*, vol. 249, p. 123483, 2024.
15. M. Singh and J. K. Chhabra, "Improved software fault prediction using new code metrics and machine learning algorithms," *Journal of Computer Languages*, vol. 78, p. 101253, 2024.
16. R. Wang et al., "An empirical assessment of different word embedding and deep learning models for bug assignment," *Journal of Systems and Software*, vol. 210, p. 111961, 2024.
17. Z. Huang et al., "On the effectiveness of developer features in code smell prioritization: A replication study," *Journal of Systems and Software*, vol. 210, p. 111968, 2024.
18. R. Suryawanshi and A. Kadam, "Software defect prediction by logistic regression with gradient descent cost computation," in *Proc. Int. Conf. Emerging Smart Computing and Informatics (ESCI)*, 2024, pp. 1–5.
19. Q. Zhang et al., "Software defect prediction using deep Q-learning network-based feature extraction," *IET Software*, 2024.
20. X. Fan, L. Lian, L. Yu, W. Zheng, and Y. Ge, "Software defect prediction method based on stable learning," *Computers, Materials & Continua*, vol. 78, no. 1, 2024.
21. Y. Qiao, L. Gong, Y. Zhao, Y. Wang, and M. Wei, "Demuvgn: Effective software defect prediction model by learning multi-view software dependency via graph neural networks," *arXiv preprint arXiv:2410.19550*, 2024.
22. M. F. Sultan, T. Karim, M. S. H. Shaon, M. Wardat, and M. S. Akter, "A combined feature embedding tools for multi-class software defect and identification," *arXiv preprint arXiv:2411.17621*, 2024.

23. R. Bibyan, S. Anand, A. Jaiswal, and A. G. Aggarwal, "Bug severity prediction using LDA and sentiment scores: A CNN approach," *Expert Systems*, vol. 41, no. 7, p. e13264, 2024.
24. F. Khan, S. Kanwal, S. Alamri, and B. Mumtaz, "Hyper-parameter optimization of classifiers using an artificial immune network and its application to software bug prediction," *IEEE Access*, vol. 8, pp. 20954–20964, 2020.
25. T. S. Mian and A. Alsaedi, "Software bug severity prediction using convolutional neural network and BiLSTM models," in *Proc. Int. Conf. Reliable Information and Communication Technology*, Cham, Switzerland: Springer, 2023, pp. 1–12.
26. A. H. Dao and C. Z. Yang, "Automated priority prediction for bug reports using comment intensiveness features and SMOTE data balancing," *International Journal of Software Engineering and Knowledge Engineering*, vol. 33, no. 3, pp. 415–433, 2023.
27. H. Hairani and D. Priyanto, "A new approach of hybrid sampling SMOTE and ENN to improve the accuracy of machine learning methods on unbalanced diabetes disease data," *International Journal of Advanced Computer Science and Applications*, vol. 14, no. 8, 2023.
28. E. Mashhadi, H. Ahmadvand, and H. Hemmati, "Method-level bug severity prediction using source code metrics and LLMs," in *Proc. IEEE 34th Int. Symp. Software Reliability Engineering (ISSRE)*, 2023, pp. 635–646.
29. F. Johnson, O. Oluwatobi, O. Folorunso, A. V. Ojumu, and A. Quadri, "Optimized ensemble machine learning model for software bugs prediction," *Innovations in Systems and Software Engineering*, vol. 19, no. 1, pp. 91–101, 2023.
30. R. S. Rao, S. Dewangan, A. Mishra, and M. Gupta, "A study of dealing class imbalance problem with machine learning methods for code smell severity detection using PCA-based feature selection technique," *Scientific Reports*, vol. 13, no. 1, p. 16245, 2023.
31. W. Yao, M. Shafiq, X. Lin, and X. Yu, "A software defect prediction method based on program semantic feature mining," *Electronics*, vol. 12, no. 7, p. 1546, 2023.
32. A. Khalid, G. Badshah, N. Ayub, M. Shiraz, and M. Ghouse, "Software defect prediction analysis using machine learning techniques," *Sustainability*, vol. 15, no. 6, p. 5517, 2023.
33. H. Kaur and A. Kaur, "SMOTE-based homogeneous prediction for aging-related bugs in cloud-oriented software," *Journal of Information & Knowledge Management*, vol. 22, no. 5, p. 2350047, 2023.
34. M. Mahalakshmi, M. P. Ramkumar, and G. S. R. Emil Selvan, "SCADA intrusion detection system using cost-sensitive machine learning and SMOTE-SVM," in *Proc. 4th Int. Conf. Advances in Computing, Communication Control and Networking (ICAC3N)*, 2022, pp. 332–337.
35. R. Mo, S. Wei, Q. Feng, and Z. Li, "An exploratory study of bug prediction at the method level," *Information and Software Technology*, vol. 144, p. 106794, 2022.
36. T. Hirsch and B. Hofer, "Using textual bug reports to predict the fault category of software bugs," *Array*, vol. 15, p. 100189, 2022.
37. D. Li, M. Liang, B. Xu, X. Yu, J. Zhou, and J. Xiang, "A cross-project aging-related bug prediction approach based on joint probability domain adaptation and k-means SMOTE," in *Proc. IEEE 21st Int. Conf. Software Quality, Reliability and Security Companion (QRS-C)*, 2021, pp. 350–358.
38. E. Mashhadi and H. Hemmati, "Applying CodeBERT for automated program repair of Java simple bugs," in *Proc. IEEE/ACM 18th Int. Conf. Mining Software Repositories (MSR)*, 2021, pp. 505–509.
39. B. S. Neysiani and S. M. Babamir, "Automatic duplicate bug report detection using information retrieval-based versus machine learning-based approaches," in *Proc. 6th Int. Conf. Web Research (ICWR)*, 2020, pp. 288–293.