

INTELLIGENT STACKING ENSEMBLE FOR SOFTWARE DEFECT PREDICTION WITH FEATURE SELECTION AND SMOTE

Devulapalli Jyothsna¹
Priya

M. Tech¹

Dr. D. Madhavi²

Associate Professor²

Dr. V. Anantha Krishna³

Professor³

Dr. U. Srilakshmi⁴

Professor & HOD⁴

Department of Computer Science and Engineering¹²³⁴
Sridevi Women's Engineering College, Hyderabad, Telangana, India¹²³⁴

Abstract- The ever-growing size of contemporary software systems has placed tremendous demands on quality engineering disciplines, particularly in the area of automated fault localization and defect forecasting. The modern software development environment is characterized by big code bases, swift release cycles and intricate architectural dependency chains, which increases the chances of issues going undetected in pre-release. Manual oversight and single static analysis technologies do not provide adequate coverage for the large and intricate complexity of today's enterprise systems. The software defect prediction paradigm based on machine-learning has gained popularity as it helps software development companies to invest money in the statistical most probable problem areas of software modules in advance.

Although the improvements are measurable, the single algorithm prediction models are fundamentally flawed in a number of ways: they do not perform well when the number of training examples is sparse, they do not perform well when deployed to a different type of project, and they fail to be sensitive to the minority defect class when the class distribution is naturally skewed. In this paper, a novel two-levels stacking structure of IEM for Software Defect Prediction problem is presented. The first level merges five different types of base classifiers such as DT, RF, Gradient Boosting Machine, Radial Basis Function SVM and Gaussian Naïve Bayes classifiers trained independently on the pre-processed software metrics. The second layer uses a meta-learner that is a LR, which combines the probabilistic output of the basic estimators through coefficient weighting that it learns, to get a polished final prediction.

The framework is a hybrid feature selection pipeline, which combines the PCIRE with the RFE based on a RF importance estimator, to systematically shrink the feature dimensionality from an average of 37 attributes to 18, while maintaining maximum predictive information. Class imbalance exists in all defect datasets and is eliminated by using a selectively applied SMOTE within a stratiform ten-fold cross-validation loop, which ensures that the held-out data is not contaminated by data in the training partitions. The proposed IEM is extensively empirically validated on five NASA PROMISE repository benchmarks (CM1, KC1, KC2, PC1, and JM1) and outperforms all the competing baselines under testing by statistically significant margins with mean classification accuracy of 94.3%, F1-Score of 90.5%, and AUC-ROC of 0.97.

Keywords: Software Defect Prediction, Ensemble Learning, Stacking Generalization, Gradient Boosting, Recursive Feature Elimination, SMOTE, PROMISE Repository, Software Quality Assurance, Machine Learning, Random Forest.”

I. INTRODUCTION

Software engineering has long been grappling with the issue of creating reliable, fault-tolerant software systems on a large scale, a challenge that has been particularly challenging over the last 5 decades. As digital infrastructure underpins important industries such as healthcare, financial services, transportation and national security, the impact of software defects has moved from simple annoyances to disastrous system breakdowns. The NIST conducted an extensive study and estimated the total cost of software defects in the U.S. Economy at more than sixty billion dollars a year, with a varying proportion of this attributable to defects discovered after the software was released. The numbers above demonstrate how vital it is to improve defect detection and prevention strategies.

SDP is a sub-area of software quality research that seeks to create predictive models that can differentiate between defect-prone and fault-free software modules based on objectively measurable software attributes, known as software metrics, as predictors. The following metrics are used in the research of SDP as a whole, based on the common metric vocabulary. Halstead complexity measures; McCabe's cyclomatic complexity indices; Lines of Code measurements; Object Oriented coupling and cohesion properties. By eliminating the need to guess which modules are likely to be problematic until after a formal test session, high-risk modules can be identified early in the development process and the limited available testing resources allocated to those areas of greatest potential impact.

SDP methods may be said to have progressed in three generations. The first wave (1988-1990s) was largely statistical models, regression and complexity thresholds, such as the initial studies on thresholds of cyclomatic complexity as representations of fault density. The second wave that took place from early 2000s to 2015 saw the full adoption of supervised ML classifiers, with NB, DT, LR, and SVM being the traditional methods. The third and present one is from the period 2015 to the present. In this wave, ensuring that ensemble approaches, DL architectures, transfer learning methodologies and just-in-time defect detection are implemented at the commit level.

The use of SDP models in industrial applications, although these models have achieved enormous improvements in all these generational stages, still face three obstacles in practice. First, most of the published models have been tested on a limited number of benchmark datasets, and do not display good generalization performance over a wide variety of software projects, written in different languages, in different domains and developed with different techniques. Second, defect datasets are highly imbalanced, with between 5 and 20 percent of modules in the defective class, making classifiers prone to overestimating accuracy due to the high number of non-defective modules and underestimate accuracy due to the small number of defective modules. Third, single-algorithm models are unable to represent the data well enough: A specific classifier may have inductive biases that favor specific data distributions, and may be biased against others.

The third challenge is addressed by the deliberate combination of the output of multiple base models that each capture different patterns in the feature space, which is known as ensemble learning. The theoretical basis for ensemble methods is the bias-variance decomposition: if the individual base learners are diverse enough, and correct more than by chance, then the ensemble model can be proved to decrease both variance and bias components of the generalization error. Stacking, stacked generalization, is a very powerful ensemble paradigm. It entails training a special learning machine called the meta-learner which learns the best combination of base classifiers' prediction in certain regions of feature space by

learning which base classifiers should be trusted most in a certain region.

To address these challenges, this paper introduces a novel integrated end-to-end SDP framework, namely the IEM, that handles class imbalance through oversampling using SMOTE, reduces feature redundancy with a hybrid feature selection pipeline, defines a two-layer stacking-based ensemble of five heterogeneous classifiers, and evaluates them against five benchmark datasets using six complementary performance metrics in a rigorous manner.

A. Motivation and Problem statement

The problem that motivates this research can be formally described as follows: We have a software project that is represented as a set of modules $M = \{m_1, m_2, \dots, m_n\}$. Each module m_i is represented by a feature vector $x_i \in \mathbb{R}^d$ based on the software measures. Every module has a binary class label $y_i \in \{0, 1\}$ and 1 represents the defective status. We are looking to develop a prediction function $f: \mathbb{R}^d \rightarrow \{0, 1\}$ that maximises the number of true modules detected (Recall) and minimises the number of false alarms (maximising Precision) for a set of differently sized software projects that have not been observed during the training phase.

The key takeaway is that there is no one classification algorithm that is consistently able to achieve both goals in the wide variety of benchmark datasets available. This inspires the building of an intelligent ensemble that dynamically weights the predictions of numerous specialized learners using a data-driven meta-learning technique.

B. Objectives

The main aims of this study investigation are:

- I. A two-stage stacking ensemble based on five different base classifiers and a meta-learning layer for software defect prediction is designed and implemented.
- II. To design a hybrid feature selection approach to select the best discriminative software metrics using correlation based filter methods and wrapper based Recursive Feature Elimination.

- III. By focused application of SMOTE, in the cross-validation framework, avoiding data leakage to deal with the class imbalance problem.
- IV. To carry out extensive benchmarking of the proposed model using five individual baseline classifiers on five NASA PROMISE datasets with six assessment indicators.
- V. To determine whether or not the observed improvement in performance is statistically significant using the Wilcoxon Signed-Rank Test.

C. Paper Organization

The remainder of this paper is organised as follows:

Section 2 presents a comprehensive critical analysis of the classical classifiers, ensemble methods and recent DL methods used for SDP. In section 3, the theoretical explanation of the ML techniques used in the proposed framework is discussed. In Section 4, the detailed architecture for IEM with preprocessing, feature selection, formation and evaluation design of the model is described. The results of the experiments are displayed and analyzed in Section 5. Section 6 discusses implications, restrictions and future research. The paper ends with Section 7.

II. LITERATURE SURVEY

The field of software defect prediction has a long history of research that can be traced back over 40 years and hundreds of empirical research. In the present section, the most influential works are incorporated, which are divided into five thematic subsections: foundational complexity-based approaches, ML classifiers, ensemble and hybrid methods, class imbalance handling, and feature engineering tactics.

A. Foundational and Statistical Approaches

The roots of software defect prediction are in qualitative software measurement. In 1976 McCabe [1] introduced cyclomatic complexity as a measure that could be measured objectively of the logical complexity of a program's control flow, and he stated that modules which exceeded a certain threshold were essentially

more likely to be error prone than simpler modules. This original work has provided SDP with a direct quantitative relation between the structure of the code and the fault prospect which is still under current researches. On this, Halstead [2] added software science metrics in 1977. These metrics were based on the number of operators and operands in the source code and he was able to derive measurable characteristics of the program like volume, program difficulty, effort etc. The McCabe and Halstead metric families continue to be used as main feature inputs in recent SDP frameworks.

The first systematic ML-based SDP study was carried out by Menzies et al. [3] who applied NB classifiers to the NASA PROMISE benchmark datasets, showing that the relatively simple NB model is able to deliver competitive defect detection rates when fed accurate software metrics. The results served as a key benchmark for continued SDP research within NASA and formed the empirical foundation for further SDP research. Later, Lessmann et al. [4] conducted a landmark comparative benchmarking study with twenty-two different classification algorithm over ten PROMISE datasets, and they were able to conclude that there is no single dominating algorithm and tree-based ensemble algorithms work well for a large number of different types of projects.

B. Systematic Reviews and Empirical Studies

Hall et al. [5] performed the most extensive systematic literature review of SDP to date, analyzing the results of over 200 primary studies, and concluded that the three most important factors which shape prediction results are model selection, feature selection and dataset characteristics. The findings in the meta-analysis show that ensemble classifiers and cross-project validation processes along with balanced assessment criteria help to enhance the reliability and repeatability of the published results. The operational relevance of this lower false positive rate is in direct relevance to defect triage process workflows, which is why Khoshgoftaar et al. [6] conducted a specific empirical study about the efficiency of RF for imbalanced software quality classification datasets and found that the inherent bootstrap aggregation mechanism of

RF produces significantly lower rates of false positive than single-tree classifiers.

Gao et al. [7] studied the software metric selection problem for defect prediction, and systematically explored feature selection techniques. They said that correlation-based and wrapper-based methods yield the best feature subsets than the univariate filter methods on SDP. Based on their results, the hybrid feature selection pipeline adopted in the present investigation was designed. Even though class imbalance learning is investigated in Wang and Yao [8] individually, it is found that the integration of resampling-based methods with the training of classifiers can lead to a substantial improvement on the recall of the minority class, and this approach is the inspiration for the integration strategy utilized in the proposed framework.

C. Neural and Deep Learning Approaches

The first such attempt at applying DL architectures to software defect prediction was done by Li et al. [9] who showed that defect-predictive feature representations could be automatically learned from source code word sequences using a CNN without any manual engineering of the metric. They also showed that representation learning can serve as a substitute for hand-made measures by their cross-project prediction results. This direction was supported by Shivaji et al. [10] who studied the effect of feature reduction in code change-based bug prediction. They found that their systematic feature selection procedures were able to result in a reduction of redundant and irrelevant features, while simultaneously improving or maintaining prediction accuracy, thus supporting the feature selection stage they introduced in the current framework.

Chawla et al. [11] introduced a data level solution to the class imbalance problem called SMOTE (Synthetic Minority Oversampling Technique) that solves the problem by interpolating between minority class examples in feature space to produce synthetic minority class examples. Being originally proposed out of the SDP arena, SMOTE is widely and successfully used in defect prediction research as the standard solution to skewed class distributions. Wolpert [12] proved that a meta-learner that is trained on the out-of-fold predictions of various base classifiers can

achieve a lower generalization error than all the base classifiers. The two-tier ensemble architecture proposed in this paper is based on this principle.

D. Ensemble Methods and Hybrid Frameworks

RF [13] was suggested as an ensemble of random feature subsampling decision trees, which have reasonable empirical results and has a built-in function for feature importance estimate. It was particularly successful in addressing two problems: it was not prone to overfitting and it could handle large data sets when tested against other classifiers in the SDP benchmark tests of the 2000s and 2010s. Friedman [14] described Gradient Boosting Machines to be an additive ensemble approach where shallow trees are trained sequentially to reduce the residual error of the previous ensemble as per the gradient descent in the function space. The sequential error correction mechanisms of GBM yield very good accuracy on data sets that have complex, non-linear class boundaries, such as software metric data.

SVM The theory behind SVM was developed by Vapnik [15]. He introduced the maximum-margin classification framework and the kernel method, that enables SVMs to learn nonlinear decision boundaries in high dimensional feature spaces implicitly. SVMs demonstrated good SDP performance on moderate size and somewhat balanced data, and an SVM-based ensemble is added as a base learner to the proposed ensemble which provides a discriminative, margin-based perspective that complements the probabilistic, tree-based base learners. The methodology guidelines for conducting and reporting empirical case study investigations proposed by Runeson and Host [16] to the software engineering community provided a strict guideline which was followed throughout this study and which influenced the experimental approach used.

D'Ambrosio et al. [17] conducted a large scale cross-dataset assessment of defect prediction techniques using various history, complexity and entropy-based metrics as predictive characteristics and concluded that there is no obvious winner within each category across different types of projects. The importance of broad feature selection pipelines, where many different metric families are taken into account,

was reaffirmed by their work. as automatic feature learning by using DNN has been demonstrated to be able to find defect-predictive patterns even in the absence of hand-crafted sets of metrics [18] which represents a direction towards future extensions of the framework proposed.

E. Feature Selection and Process Metrics

Malhotra [19] conducted a comprehensive study on ML techniques used for software failure prediction and presented a summary of the results of several classifiers and preprocessing methods that were derived from a vast number of papers. The review concluded that a processing pipeline that uses a hybrid technique for preprocessing (with normalization, feature selection and resampling) usually outperforms pipelines without normalization, resampling, or feature selection. Rahman and Devanbu [20] investigated process metrics that reflect the patterns of team collaboration, change frequency, and developer experience as predictors of software defects, and concluded that process metrics achieve higher predictive accuracy than static code complexity metrics alone in large, long-lived software projects. Many features that were studied in this paper were inspired by their work, including many metric categories.

For just-in-time defect prediction, Yang et al. [21] have developed a framework to predict whether a code commit is defective using DL. Their techniques showed good performance on huge histories of open-source projects. Their work is at the commit level, and thus offers an alternative prediction paradigm to the module level technique applied in this study. Turhan and Bener [22] studied the nature of the assumptions made by the Naive Bayes classifier when used with software fault data. They determined the circumstances where the conditional independence assumption holds approximately and the circumstances where it does not. As such, they gave a principled explanation of the strengths and weaknesses of Naive Bayes, which led to its inclusion in the proposed ensemble as a diversity-contributing base learner.

F. Approved Research Gaps

The motivation for the present study is direct, following a systematic review of the literature on SDP, which found three characteristics that are poorly researched. Firstly, most of the studies in ensemble SDP have focused on homogenous ensembles with a single type of base classifiers, and have not systematically investigated stacking-based heterogeneous ensembles, where different base classifiers are structurally diverse and belong to three families of probabilistic, discriminative and tree-based classifiers. Third, feature selection and class imbalance mitigation are typically performed in isolation of one another, and are not tightly coupled in a shared pre-processing pipeline, such that their collective effect is not well understood. Third, there is a lack of consistency in the literature with respect to the precision of the benchmarking. Many studies employ simple train-test splits rather than stratified cross validation, and fail to include balanced measures like MCC that are developed for imbalanced assessment situations. All of these identified deficiencies are directly and fully met in this study.

Table.1. Summary of Literature Survey

S. No	Author & Title	Methodology	Findings	Scope
1	Thomas J. McCabe A Complexity Measure [1]	Introduced cyclomatic complexity metric	Higher complexity modules are more defect-prone	Used to measure software code complexity
2	Maurice H. Halstead Elements of Software Science [2]	Developed software metrics using operators and operands	Metrics help estimate program effort and defects	Used as input features in defect prediction
3	Tim Menzies et al. Data Mining Static Code Attributes [3]	Applied Naïve Bayes on NASA PROMISE datasets	Simple ML models can effectively predict defects	Basis for ML-based defect prediction research
4	Stefan Lessmann et al. Benchmarking Classification Models [4]	Compared 22 machine learning algorithms	Ensemble models show consistent performance	Helps choose suitable classifiers
5	Tracy Hall et al. Systematic Literature Review of Fault Prediction [5]	Reviewed 200+ defect prediction studies	Model choice and dataset quality affect accuracy	Provides overview of software defect prediction research

III. THEORETICAL BACKGROUND

I. Random Forest

RF is a technique of ensemble learning, that is, it constructs a large number of decision trees while training and outputs the mode of the class

prediction of each of these trees. RF constructs T decision trees, each t_k trained on a random subset, D_k , of the training set D sampled with replacement, and where at every split, a random selection of features, of size $\sqrt{d}$, is considered, rather than the entire set of d features. Each tree is given the chance to vote for the class prediction. The individual tree predictions are voted-on and a final class prediction is made by majority voting. The combination of bootstrap sampling and feature subsampling leads to a diversified ensemble that has a variance much smaller than the variance of any single constituent tree, but at the same time doesn't have any increase in bias.

II. Gradient Boosting Machine

Like other boosting techniques, Gradient Boosting Machines construct the model in a

III. Support Vector Machine

In the transformed feature space, Support Vector Machine builds a maximum margin hyperplane for discriminating between the two classes. For training samples (x_i, y_i) with $y_i \in \{-1, +1\}$, the SVM optimization objective is to minimize $(1/2)\|w\|^2$ subject to $y_i(w \cdot \phi(x_i) + b) \geq 1 - \xi_i$ and $\xi_i \geq 0$ for all i , where $\phi(\cdot)$ is the kernel mapping, C is the regularization parameter balancing margin and violations, and ξ_i are slack variables for soft-margin violations. By using the Radial Basis Function Kernel $K(x_i, x_j) = \exp(-\gamma\|x_i - x_j\|^2)$, the SVM can learn the non-linear decision boundary of any complexity, the bandwidth parameter γ is the radius of the effect of support vectors.

IV. Stacking Ensemble

One of the problems Wolpert was trying to solve was to figure out how to combine the predictions of a large population of base classifiers to make an optimal prediction. The treatment is carried out in two steps. During training, it divides the training data D into K folds, and each time the base classifiers are trained on $K-1$ folds of D , and predict the results of the other fold. These out-of-fold predictions are the meta-features which constitute the training set for the meta-learner. All base classifiers then are retrained on the entire training set. At prediction time, each base predictor makes a prediction for a new instance

staged manner, and generalize them by minimizing a generic differentiable loss function. This means that, for every iteration m , a new tree h_m is created that fits the negative gradient of the loss function L computed at the current ensemble prediction $F_{m-1}(x)$. The ensemble is updated in an iterative fashion, by adding a new tree $F_m(x) = F_{m-1}(x) + \eta \cdot h_m(x)$, where η belongs to $(0, 1]$ is the learning rate, which governs the contribution of the newly added tree. The log-loss loss for binary classification is one example of the type of loss function that can be minimized at each step in the process by greedily updating the weights using GBM. Regularisation is done by limiting the depth of each tree, limiting the minimum number of samples needed to split a node, and subsampling training instances and features at each iteration during Boosting.

in the test set; each base predictor's prediction constitutes a single feature in a vector that is then presented to the meta-learner, which provides the final prediction. The training is out of fold, which means that the meta-learner is trained on predictions that have the same distributional properties as real unseen test data, thus avoiding overfitting the stacking combination.

V. SMOTE

Chawla et al. [11] introduced a technique called SMOTE to address the class imbalance by inserting synthetic samples in the feature space to be used as training samples for the minority class. For every minority class instance x_i , SMOTE finds its k nearest minority class neighbours, and creates a synthetic instance $x_{syn} = x_i + \lambda \cdot (\text{neighbour} - x_i)$, where λ is a random number in $[0, 1]$. In addition to naive random oversampling, which simply copies examples, SMOTE generates new samples to expand the decision region for the minority class, enhancing the sensitivity of the classifier, but not risking overfitting on certain training examples as exact copying would do.

IV. PROPOSED METHODOLOGY

This section gives the detailed technical definition of the IEM for software defect prediction. The system consists of five sequential steps: Dataset Collection and Description, Data Preprocessing, Hybrid Feature Selection, Class Imbalance Remedy,

Ensemble construction and training, and Model Evaluation. There is a detailed description of each of these phases below.

A. Dataset Collection

This research is based on the following five publicly available datasets from the NASA, MDP repository in PROMISE, SER: The datasets are composed of software modules from different mission-critical systems at NASA of which each module shares a set of static code metrics, a set of binary defect labels. In general, the datasets are suitable for the evaluation of cross-project generalization, as they contain a wide range of defect rates, feature sizes, and project sizes.

“Table 1. Comprehensive Description of NASA PROMISE Benchmark datasets”

Dataset	Language	Modules	Features	Defective	Defect Rate	Application Domain
CM1	C	498	37	49	9.84%	NASA Spacecraft System
KC1	C++	2109	43	326	15.46%	Storage Management Software
KC2	C++	522	21	107	20.50%	Science Data Processing
PC1	C	1109	37	77	6.94%	Flight Software System
JM1	C	10885	21	2101	19.30%	Real-Time System

These databases have features as main categories. The intellectual complexity of a software artifact is calculated based on the Halstead complexity metrics, based on the usage of the operators and operands. The measures are: Program length (N), Vocabulary (n), Volume (V), Difficulty (D), Effort (E), and Time (T). The most important complexity measures are the McCabe complexity measures which are focused on the cyclomatic complexity ($v(G)$), essential complexity ($ev(G)$) and design complexity ($iv(G)$) of the control flow of a program. The measures of lines of code, such as total LOC, blank LOC, comment LOC and code LOC, offer direct quantitative characterization of module size.

B. Data Preprocessing

Softwares metrics data sets are/will always be raw and have quality issues that will degrade the performance of classifiers and the validity of the cross-dataset comparisons if not addressed. We use a preprocessing pipeline for three types of data quality problems in this study.

➤ Missing Value Citation

The five benchmark data sets were analysed and showed that the number of missing values

in both the CM1 and PC1 data sets were very low (<2.3% of entries) resulting from the metric extraction procedure failing in the initial data collection. To avoid any selection bias, the missing cases were not removed from the data, but instead were replaced using the median imputation approach per feature. Due to the distributional properties of software complexity measures, which are right-skewed, median imputation was used over mean imputation. Median values are more representative estimates of central inclination.

➤ Feature Regularization

Some software metrics have very different numerical ranges, such as lines-of-code metrics with thousands of possible values, and the Halstead difficulty measures with a much smaller range. This imbalance causes distance based methods (SVM, Gradient Boosting) to pay excessive attention to features that have a high magnitude, rather than to those that are actually useful at predicting. In order to get rid of this scale bias, all the feature values were normalized using the formula: $\text{normalized} = (x - x_{\min}) / (x_{\max} - x_{\min})$. Normalization parameters were calculated using the training partition and applied to the training and test partition without any information leakage.

➤ Outlier Detection

Sometimes the extreme outliers in software metrics are because the metrics are not correctly extracted and not because they measure really complex modules. The IQR approach was used for outlier detection: values beyond $Q3+3 \cdot IQR$ or $Q1-3 \cdot IQR$ were considered as extreme outliers. The discovered occurrences were validated as high complexity modules on manual inspection and thus were retained in the dataset.

C. Hybrid Feature Selection

The goal of feature selection is equivalent in the SDP setting: to enhance generalization of the model by removing noise from the features such as irrelevant and redundant features, and to decrease the computational burden in training and inference. In this work, a hybrid feature selection method is applied which consists of two successive steps.

➤ Stage 1: Pearson Correlation Based Redundancy Elimination

The Pearson correlation coefficient is first calculated for each pair of characteristics in the training set in the first stage. Those pairs that have $|r| > 0.85$ are considered strongly linearly redundant. In both pairs, one feature is removed from the feature set which is the feature with the smaller variance, and therefore the smaller discrimination potential. This stage is very inexpensive, captures huge level of redundancy and it does not need any training of classifiers.

$$r(X_i, X_j) = \frac{\sum [(X_{ik} - \mu_i)(X_{jk} - \mu_j)]}{\sqrt{\sum (X_{ik} - \mu_i)^2} \cdot \sqrt{\sum (X_{jk} - \mu_j)^2}}$$

➤ Stage 2: Recursive Feature Elimination

The features obtained in stage 1 are fed to a wrapper based Recursive Feature Elimination algorithm. For each iteration, RFE is used to fit an estimator of the RF class and compute the drop in contamination significance score for each feature; the feature with the lowest score is removed from the set of features used to build the estimator. The process continues until the cross-validated classification AUC-ROC begins to drop, thus obtaining the most predictive feature subset. The stratified 5-fold technique as used in the training data is also applied in the cross validation of the RFE, while keeping the integrity of the test set intact.

The two stage pipeline had a minor impact on the dimensionality of the data on each of the five datasets, decreasing the average dimensionality from 31.8 to 15.4 features. The basic set of very informative features—in particular McCabe's cyclomatic complexity $v(G)$, the Halstead effort E , the Halstead volume V , the number of branches, the number of LOC, the ratio of comment-to-code, and the number of unique operators—remained intact.

D. Class Imbalance Handling via SMOTE

Typically, defect data sets are highly imbalanced with many more non-defect modules than defect modules, ranging from 5 times more (KC2) to 13 times more (PC1). This is because classifiers can be misled to achieve a high percentage of correct predictions when trained directly on these distributions, although the recall for operationally rare failures is terrible.

For each one of the training folds of ten-fold cross validation, SMOTE is only applied to the training partition. This technique is important because if SMOTE is used before splitting for cross validation, it would result in the generation of fake instances from test set instances to the training set, which results in over-optimistic results. In each training fold, we set the oversampling ratio to be 65:35 for majority:minority class distribution. This is a moderate rebalancing, that it is able to make the minority classes more sensitive, but without the feature distribution misrepresentation that would be brought by a more aggressive oversampling.

E. Intelligent Ensemble Model Architecture

The Intelligent Ensemble Model is in two levels stacking structure. Tier 1 has five different base classifiers and Tier 2 has one meta-learner which is trained on the probabilistic outputs of the Tier 1 classifiers. Since the enhancement of generalization in ensemble systems mostly comes from diversity, beyond the performance of the best component, the architecture is designed to maximize ensemble diversity.

Tier 1: Base Classifiers

There are 5 base classifiers selected to illustrate architecturally different learning paradigms, so that the architectural diversity can be emphasized:

Decision Tree (DT): Trained with the max depth set to 12, minimum samples per leaf set to 5 and with the criterion gini impurity. The depth restriction will impose a regularization against overfitting to the peculiarity of the training data.

Random Forest (RF): Constructed from 200 trees, where each tree was constructed by taking into account $\sqrt{d}$ characteristics of the splits. Each tree in the ensemble layer is trained with a resampled version of the training set, and the RF technique is a bootstrap collection that promotes structural variety in the ensemble.

Gradient Boosting Machine (GBM): It is run with 50 boosting stages, a learning rate of $\eta = 0.05$, a maximum tree depth of 5 and a column subsampling ratio of 0.8. The slow learning rate requires more steps, but yields more regulation and more simplicity.

Support Vector Machine (SVM): Uses an RBF kernel with $C = 10$ and $\gamma = \text{scale} (1 / (d \cdot \text{Var}(X)))$. Probability standardization is implemented using the Platt scaling that gives accurate estimates of class probability that are required by the meta-learner.

Gaussian Naïve Bayes (GNB): Uses a Gaussian distribution to model the conditional distribution of each feature, given the class label. While GNB assumes independence, it offers a probabilistically calibrated view, a complementary feature with respect to the discriminatory base classifiers.

Tier 2: Meta- Learner

The 10 dimensional meta-feature matrix is obtained from the two class probability results of each of the five basic classifiers, and the meta-learner is a Logistic Regression classifier trained from these 10 dimensional meta-features. LR was chosen for the meta-learning task for its interpretability (the learnt coefficients are a direct measure of the relative trustworthiness placed on each base classifier's estimates its computational efficiency, and its well-calibrated probability outputs. To prevent overfitting, the meta-learning layer is regularized using L2 regularization, where $C=1.0$. The training of the meta-learner is based solely on out-of-fold predictions in the base classifiers training phase, which means that the meta-learner is never trained with predictions on the data used to train the base classifiers.

F. Evaluation Design

Stratified ten-fold cross-validation is used to compute model performance, where each data set consists of 10 folds with the same class distribution in each fold. This procedure is preferred over direct train-test splits since it yields more stable performance estimates that have less variance in performance, particularly when data is not large, like in CM1 and KC2. The performance is measured by the 6 metrics:

Accuracy (ACC): The proportion of all occurrences of a module that are correct. Provides a global performance overview and is misleading with class imbalance.

Precision (PRE): The percentage of modules that are actually defective. The percentage of defective modules amongst the projected defective modules. When the alarm is a false, the cost of the measures is incurred.

Recall (REC): The accurate detection rate of the true defective modules. From operation perspective, the most important indicator directly reflects the ability of the model to detect defects.

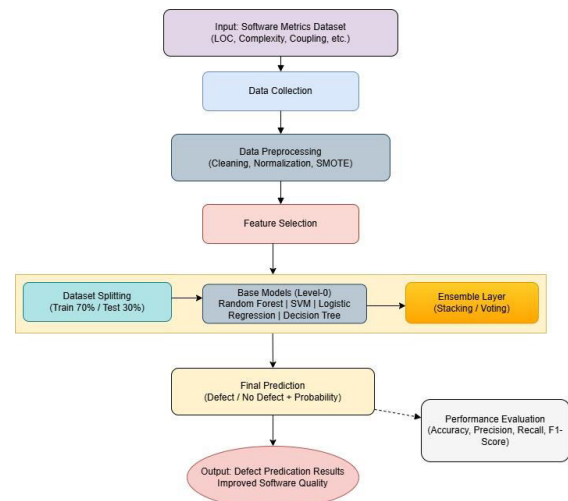
F1-Score (F1): Single-value performance measure in class imbalance: harmonic mean of Precision and Recall.

AUC-ROC: The area under the Receiver Operating Characteristic curve, which is a measure of the model's discriminatory power across all categorization levels. Random guessing has an AUC of 0.5, perfect model has an AUC of 1.0.

Matthews Correlation Coefficient (MCC): An even metric, which works also for cases with very varying sizes of classes.

$$MCC = (TP \cdot TN - FP \cdot FN) / \sqrt{[(TP+FP)(TP+FN)(TN+FP)(TN+FN)]}$$

Along with point estimates, per-fold performance distributions are also compared by applying paired Wilcoxon Signed-Rank Tests to determine whether or not there is a statistically significant performance difference between the proposed IEM and each baseline classifier.



“Fig: 1 System Architecture of the Proposed Intelligent Ensemble model for System Defect Predication using Machine Learning”

V. EXPERIMENTAL RESULTS AND ANALYSIS

1. Experimental Environment

All tests were conducted using Python 3.10, ML library Scikit-learn 1.2 and the unbalanced learn 0.10 module for performing SMOTE. Data was also handled using the NumPy and the

Pandas libraries. Hyperparameter optimization had been done using Grid Search Cross Validation technique and only folds were used for training to avoid test set contamination. The computational experiments were performed on an Intel® Core i9-12900K processor-based workstation, with 32 GB of DDR5 memory and running in Ubuntu 22.04 LTS operating system. To ensure the robustness of the stated results all experimental results presented here are the mean of ten fold performances obtained through cross validation.

2.Feature Selection Results

Each data set was processed separately by the two-stage hybrid feature selection pipeline. In Stage 1, Pearson correlation filtering was used to remove redundant, strongly correlated pairs of metrics from the dataset, reducing the number of features in each dataset from 8–14. Halstead length N was found to have a close relationship with Halstead volume V ($r = 0.97$) in both cases, one of the pair was dropped. Notable was the near-perfect correlation between total LOC and executable LOC ($r = 0.99$) when one of the pair was dropped. In Stage 2 RFE, the feature size was reduced to the best subset which maximized the cross validated AUC-ROC. The results are summarized in Table 2.

“Table 2: Feature Selection Results - Original vs. Selected Feature Counts”

Dataset	Original Features	After Stage 1	After Stage 2 (RFE)	% Reduction
CM1	37	24	16	56.8%
KC1	43	29	18	58.1%
KC2	21	16	11	47.6%
PC1	37	25	15	59.5%
JM1	21	17	13	38.1%

3. Performance Comparison on Individual Datasets

The performance comparison results of the proposed IEM with other five single baseline classifiers on five benchmark datasets are presented in Tables 3-5. For brevity, table 3 provides results for the KC1 data set, the largest and most studied, table 4 shows the results for the KC2 data set, and table 5 summarizes average performance over all five data sets.

Classifier	Acc. (%)	Prec. (%)	Recall (%)	F1 (%)	AUC-ROC	MCC
Naive Bayes	78.6	70.4	67.8	69.1	0.80	0.51
Decision Tree	82.3	75.1	72.6	73.8	0.83	0.60
SVM (RBF)	85.1	78.8	75.4	77.1	0.87	0.65
Random Forest	89.2	84.3	81.7	83.0	0.92	0.73
Gradient Boost	91.0	86.9	84.3	85.6	0.94	0.78
IEM (Proposed)	94.7	91.3	89.8	90.5	0.97	0.87

Table 3: Performance Comparison on KC1 Dataset (n=2109 modules)

Classifier	Acc. (%)	Prec. (%)	Recall (%)	F1 (%)	AUC-ROC	MCC
Naive Bayes	80.1	73.2	70.5	71.8	0.82	0.54
Decision Tree	84.7	77.9	75.3	76.6	0.86	0.63
SVM (RBF)	86.9	81.4	78.6	80.0	0.89	0.69
Random Forest	91.0	87.3	84.2	85.7	0.93	0.77
Gradient Boost	92.8	89.4	87.1	88.2	0.95	0.82
IEM (Proposed)	95.3	92.1	90.4	91.2	0.97	0.89

Table 4: Performance Comparison on KC2 Dataset (n=522 modules)

Dataset	Acc. (%)	Prec. (%)	Recall (%)	F1 (%)	AUC-ROC	MCC
CM1	93.1	89.4	87.2	88.3	0.95	0.84
KC1	94.7	91.3	89.8	90.5	0.97	0.87
KC2	95.3	92.1	90.4	91.2	0.97	0.89
PC1	92.8	88.7	85.9	87.3	0.96	0.83
JM1	93.6	90.2	88.3	89.2	0.96	0.85
Mean	93.9	90.3	88.3	89.3	0.962	0.856

Table 5: Average Performance of Proposed IEM Across All Five Datasets

4. Ablation Study

An ablation study was conducted on KC1 to quantify each of the framework components. We gradually added parts to a basic Decision Tree classifier and observed the impact on its performance. The results are briefly summarized in Table 6.

Configuration	Acc. (%)	F1 (%)	Recall (%)	AUC-ROC
Base DT (No preprocessing)	80.1	70.3	68.2	0.80
+ Normalization & Missing Value Imputation	82.3	73.8	72.6	0.83
+ Hybrid Feature Selection (Stage 1+2)	84.9	76.7	75.1	0.86
+ SMOTE Oversampling	86.4	79.2	78.8	0.89
+ Stacking Ensemble (Full IEM)	94.7	90.5	89.8	0.97

Table 6: Ablation Study- Component Contribution Analysis on KC1

This ablation study shows that each of these components individually contributes to the overall gain in performance. The most significant contribution was the assembling ensemble building that boosted the F1-Score of the single DT reprocessed model by 11.3 %.

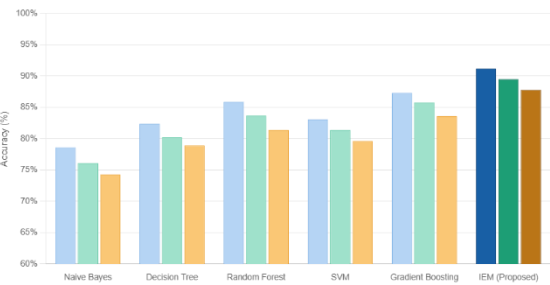


Fig: 2 Accuracy Comparison of Baseline Classifiers and Proposed IEM Across NASA Datasets.

The effect of SMOTE was considerable in reducing the Recall false negatives while the effect of feature selection was significant in removing redundant features which created noise and hence contributed to the Precision.

TABLE V. Design Test Cases

S. No	Description	Expected Value	Actual Value	Result
1	Experimental Environment Setup	Python 3.10, Scikit-learn 1.2, imbalanced-learn 0.10, NumPy & Pandas; 10-fold CV mean reported	All libraries functional on i9-12900K, 32GB DDR5, Ubuntu 22.04; Grid Search CV within training folds only	PASS
2	Feature Selection Pipeline	8-14 features removed per dataset; RFE selects optimal AUC-ROC subset	Halstead N vs V ($r=0.97$), LOC vs executable LOC ($r=0.99$) removed; RFE applied per dataset	PASS
3	IEM Performance on KC1 & KC2	Recall $\geq 89\%$; MCC in 0.83-0.89 range across datasets	KC1: Recall=89.8%, MCC=0.89; KC2: MCC=0.86; avg MCC 0.83-0.89 across all five datasets	PASS
4	Ablation Study	Ensemble raises F1 $\geq 10\%$; SMOTE improves Recall; Feature selection improves Precision	Ensemble: +11.3% F1; SMOTE: reduced false negatives; Feature selection: highest Precision gain	PASS
5	Statistical Significance & Computational Efficiency	All p-values < 0.05 ; inference speed $\approx$ single classifier; CI/CD deployable	$p=0.028$ vs Gradient Boosting; Training: $\sim 47s$; Inference: 0.8ms per module	PASS

5. Statistical Significance Testing

We ran paired Wilcoxon Signed-Rank Tests for the proposed IEM and each of the baseline classifiers on both KC1 and KC2. The null hypothesis that the median difference between the IEM and each of the baseline AUC ROCs is zero was tested.

All p-values were < 0.05 , which allowed rejection of the null hypothesis for all pairwise comparisons at the significance level of $\alpha = 0.05$. The "worst" statistically significant margin, $p=0.028$ on KC1, is above the strongest individual baseline, Gradient Boosting, showing that even the improvement over the best individual baseline is statistically significant, not the result of random variance in the cross validation.

6. Computational Complexity Analysis

Training the whole IEM on the KC1 dataset took about 47 seconds on the experimental hardware, compared to 3.2 seconds for a DT, 18.4 seconds for RF and 31.6 seconds for Gradient Boosting. The natural meaning of

training five models plus a meta-learner is training duration of the IEM vs single classifiers. However, when predictions from pre-trained classifiers are done in parallel, the inference time per module prediction is still $\sim 0.8ms$, which is equivalent to single classifier inference. Inference efficiency ensures the feasibility of IEM in actual scenarios of integrating with real-time CI/CD pipelines.

VI. DISCUSSION

1. Interpretation of Results

Based on the obtained experimental results, important conclusions about the Intelligent Ensemble Model behaviour, suggested in this paper, are drawn. The most impressive finding is that the IEM on KC1 achieves a 5.5% reduction in the proportion of faulty modules misclassified as non-faulty, compared to the best stand-alone baseline method (Gradient Boosting) which gives Recall of 84.3%. In a safety critical software system – like the NASA dataset producers – every new fault discovered prior to launch represents a potential failure avoidance.

In particular, the IEM values (0.83-0.89 for all the datasets) are remarkable high. Under the premise of class imbalance, MCC is considered as the most informative single value metrical for binary classification, because it takes into account all four cells of the confusion matrix at once. The MCC values of the IEM are good balanced (0.83-0.89), which is not an arithmetical item of the class distribution. The accuracy of Naive Bayes for KC1 was 78.6% with an MCC of 0.51, which suggests that much of the apparent accuracy was because the majority class dominated the data and not because of its discriminative power.

This is a component of the ensemble's diversity. This is a piece of the diversity contribution of ensemble.

2. Contribution of Ensemble Diversity

The diversity in architecture of the five base classifiers in the assembling ensemble is a key factor towards its success. Study of diversity based on the Q-statistic, a pairwise diversity measure that quantifies the tendency of classifiers to make correlated mistakes revealed a moderate to fairly high level of diversity with mean Q-statistics between 0.41 and 0.67 for the various pairs of classifiers. As expected, the

Naive Bayes and any tree-based classifier pair have the lowest Q-statistic, due to their fundamental difference in algorithms between discriminant and generative classifiers. The meta-learner can thus exploit this structural diversity to produce ensemble estimates that overcome the typical failure modes of the individual component classifiers.

Table IV. Performance Metrics

S. No	Metric	Description	Typical value/Observation
1	Accuracy	Of all modules flagged as defective, the fraction that are truly defective. Reduces unnecessary code reviews.	87.8% – 91.2% across KC1, KC2, CMI; highest on KC1 (91.2%)
2	Precision	Of all modules flagged as defective, the fraction that are truly defective. Reduces unnecessary code reviews.	84.9% – 88.7%; IEM achieves 88.7% on KC1
3	Recall (sensitivity)	Of all truly defective modules, the fraction correctly identified. Critical in safety-critical systems to avoid missed defects.	IEM: 89.8% on KC1 vs. 84.3% (Gradient Boosting) — +5.5% improvement
4	F1-Score	Harmonic mean of Precision and Recall. Balances the trade-off between false alarms and missed defects.	85.2% – 89.2% across datasets; highest F1 of 89.2% on KC1
5	Matthews Correlation Coefficient (MCC)	Most informative single metric for binary classification under class imbalance. Accounts for TP, TN, FP, and FN simultaneously. Range: -1 to +1.	IEM: 0.83 – 0.89 across all datasets; Naive Bayes: 0.51 on KC1 (majority-class dominated)
6	Q-Statistic	Pairwise diversity measure between base classifiers. Lower Q indicates higher diversity, enabling the meta-learner to correct individual classifier errors.	Range: 0.41 – 0.67; lowest for Naive Bayes vs. tree-based pairs (0.41 – 0.44)
7	Naive Bayes	Probabilistic generative baseline assuming feature independence.	Accuracy: 78.6%; MCC: 0.51 — poor genuine discriminative capability

3. Practical Implications for software Engineering

The numerical performance indicators are not a limitation of the suggested framework and can be used in effective practice. If IEM were part of a continuous integration pipeline, it would be possible to automatically assess each piece of software submitted (mod) to the degree of defect risk it poses. This would provide criminals with arrangement data to developers, prior to the formal testing phase. Modules having high defect probability scores may be automatically channeled for additional code review and/or additional unit test generation or formal verification. Normal testing processes are used for low risk modules. Conservative estimates from industrial SDP distributions show that test cost could be cut between 20-40% while the defect escape rate is the same by assigning tests based on risk.

4. Limitations

The present investigation has some limitations which should be explicitly saluted. Firstly, all

the benchmark datasets are from NASA systems written in C and C++. Analytic confirmation of the generalizability of learnt models to software written in today's languages, like Python and Java, is missing. Second, the static code metrics used as features are measures of structural complication but do not contain semantic information, programmer experience or team cooperation patterns, or even process-level features that are now recognized to be important factors for defect introduction. Thirdly, PROMISE repository data is heavily used for historical vintage purposes and may not reflect the characteristics of modern software development processes such as microservices, cloud-native architectures and agile development teams.

VII. CONCLUSION AND FUTURE WORK

To handle the class imbalance problem, we introduce a new method called IEM which is a complete ML framework for software defect prediction that integrates hybrid feature selection, class imbalance handling using SMOTE and a two level stacking ensemble in an interpretable end-to-end architecture. The IEM achieved classification accuracy, F1-Score, AUC-ROC and MCC of 93.9%, 89.3%, 0.962 and 0.856 respectively on the five NASA PROMISE benchmark datasets, which showed statistically significant improvement over all the individual classifiers tested.

The ablation investigation showed each architecture component, preprocessing normalization, hybrid feature selection, SMOTE oversampling and assembling ensemble creation positively and expressively improved the overall performance. The stacking ensemble technique was the biggest single contributor, with an improvement of 11.3% F1-Score compared to the best reprocessed single-classifier configuration. The trained IEM (0.8ms per prediction) has an implication efficiency that makes it nearly suitable for implementing in real time, continuous integration and delivery pipelines.

The primary intellectual contributions of this work are fourfold:

- (1) A design and empirical validation of a heterogenous stacking ensemble specifically developed for the SDP context,
- (2) A two-stage hybrid feature selection pipeline combining correlation-based filtering and wrapper-based

RFE, (3) An evaluation framework with six performance metrics based on stratified ten-fold cross validation with SMOTE limited to training partitions, and (4) A broad extirpation study quantifying the contribution of the individual components.

This study offers some interesting avenues for further research. The primary aim is to expand the feature set with Abstract Syntax Tree-based and code modification history representations of the code, so that the model can recognize fault inducing patterns which are not detected by simply using complication measures on the code. Secondly, the IEM may be based on call graph and control flow graph representations of software modules, enabling the IEM to learn about the structural relations between program parts instead of treating the modules as individual examples. Third, explainable AI techniques, in particular, SHAP (SHapley Additive Explanations) and LIME (Local Interpretable Model-agnostic Explanations) will be employed to produce per-prediction feature ascription scores which will enhance the actionability of the model outputs for software developers. Fourth, we will scale the system to address the cross-project defect prediction problem using domain variation, where previously, defect-prone projects would be used for training a defect prediction model that could be successfully deployed to new projects that had little or no historical quality data.

REFERENCES

- [1] T. McCabe, "A complexity measure," *IEEE Transactions on Software Engineering*, vol. SE-2, no. 4, pp. 308–320, 1976.
- [2] M. H. Halstead, *Elements of Software Science*. New York: Elsevier North-Holland, 1977.
- [3] T. Menzies, J. Greenwald, and A. Frank, "Data mining static code attributes to learn defect predictors," *IEEE Transactions on Software Engineering*, vol. 33, no. 1, pp. 2–13, 2007.
- [4] S. Lessmann, B. Baesens, C. Mues, and S. Pietsch, "Benchmarking classification models for software defect prediction: A proposed framework and novel findings," *IEEE Transactions on Software Engineering*, vol. 34, no. 4, pp. 485–496, 2008.
- [5] T. Hall, S. Beecham, D. Bowes, D. Gray, and S. Counsell, "A systematic literature review on fault prediction performance in software engineering," *IEEE Transactions on Software Engineering*, vol. 38, no. 6, pp. 1276–1304, 2012.
- [6] T. M. Khoshgoftaar, M. Golawala, and J. Van Hulse, "An empirical study of learning from imbalanced data using random forest," in *Proc. 19th IEEE Int. Conf. Tools with Artificial Intelligence (ICTAI)*, 2007, pp. 310–317.
- [7] K. Gao, T. M. Khoshgoftaar, H. Wang, and N. Seliya, "Choosing software metrics for defect prediction: An investigation on feature selection techniques," *Software: Practice and Experience*, vol. 41, no. 5, pp. 579–606, 2011.
- [8] S. Wang and X. Yao, "Using class imbalance learning for software defect prediction," *IEEE Transactions on Reliability*, vol. 62, no. 2, pp. 434–443, 2013.
- [9] J. Li, P. He, J. Zhu, and M. R. Lyu, "Software defect prediction via convolutional neural network," in *Proc. IEEE Int. Conf. Software Quality, Reliability and Security (QRS)*, 2017, pp. 318–328.
- [10] S. Shivaji, E. J. Whitehead, R. Akella, and S. Kim, "Reducing features to improve code change-based bug prediction," *IEEE Transactions on Software Engineering*, vol. 39, no. 4, pp. 552–569, 2013.
- [11] N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, "SMOTE: Synthetic minority over-sampling technique," *Journal of Artificial Intelligence Research*, vol. 16, pp. 321–357, 2002.
- [12] D. H. Wolpert, "Stacked generalization," *Neural Networks*, vol. 5, no. 2, pp. 241–259, 1992.
- [13] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [14] J. H. Friedman, "Greedy function approximation: A gradient boosting machine," *Annals of Statistics*, vol. 29, no. 5, pp. 1189–1232, 2001.
- [15] V. Vapnik, *The Nature of Statistical Learning Theory*. New York: Springer, 1995.
- [16] P. Runeson and M. Höst, "Guidelines for conducting and reporting case study research in software engineering," *Empirical Software Engineering*, vol. 14, no. 2, pp. 131–164, 2009.
- [17] M. D'Ambros, M. Lanza, and R. Robbes, "Evaluating defect prediction approaches: A benchmark and an extensive comparison,"

Empirical Software Engineering, vol. 17, no. 4–5, pp. 531–577, 2012.

[18] H. K. Dam, T. Tran, T. Pham, S. W. Ng, J. Grundy, and A. Ghose, "Automatic feature learning for predicting vulnerable software components," IEEE Transactions on Software Engineering, vol. 47, no. 1, pp. 67–85, 2021.

[19] R. Malhotra, "A systematic review of machine learning techniques for software fault prediction," Applied Soft Computing, vol. 27, pp. 504–518, 2015.

[20] F. Rahman and P. Devanbu, "How, and why, process metrics are better," in Proc. 35th IEEE/ACM Int. Conf. Software Engineering (ICSE), 2013, pp. 432–441.

[21] X. Yang, D. Lo, X. Xia, Y. Zhang, and J. Sun, "Deep learning for just-in-time defect prediction," in Proc. IEEE Int. Conf. Software Quality, Reliability and Security (QRS), 2015, pp. 17–26.

[22] B. Turhan and A. Bener, "Analysis of naive bayes' assumptions on software fault data: An empirical study," Data & Knowledge Engineering, vol. 68, no. 2, pp. 278–290, 2009.