

CLUSTER-GUIDED ROBUST HYBRID LEARNING FOR IMBALANCED SOFTWARE DEFECT PREDICTION

DR SAMBASIVARAO BARAGADA¹, P.SUNITHA DEVI², BAGAM LAXMAIAH³, DR V
RAMALATHA⁴, JAYAVELU S⁵, PUPPALA RAMYA⁶, ELANGO VAN MUNIYANDY⁷ DR.
SURESH KUMAR PITTALA⁸

¹ Associate Professor of Computer Science, Department of Computer Science, M.V.S Government Arts and Science College, Mahabubnagar, Telangana - 509001.

² Department of Computer Science & Engineering
G.Narayanamma Institute of Technology and Science, Hyderabad, India.

³ Associate Professor, Department of CSE, CMR TECHNICAL CAMPUS
Medchal Road, Kandlakoya, Telangana, India - 501401

⁴ Associate Professor-Selection Grade, Department of Mathematics, Presidency University, Bangalore

⁵ Department of Mechanical Engineering, Vel Tech Rangarajan Dr Sagunthala R&D Institute of Science and Technology, Avadi 600062, India.

⁶ Assistant Professor, Department of Computer Science and Engineering, Koneru Lakshmaiah Education Foundation, Green Fields, Vaddeswaram, Guntur Dist., A.P., India.

⁷ Department of Biosciences, Saveetha School of Engineering. Saveetha Institute of Medical and Technical Sciences, Chennai - 602 105.

⁸ Professor, Dept of ECE, R.V.R & J.C College of Engineering, Guntur - 522019

E-mail: shivaraao.bs@gmail.com, sunitha@gnits.ac.in, blaxmanphd@gmail.com, v.ramalatha@gmail.com, sjayavelu@veltech.edu.in, mothy274@kluniversity.in, muniyandy.e@gmail.com, dr.sureshkumarpittala@gmail.com

ABSTRACT

Software defect prediction datasets typically exhibit missing values, noisy metrics, severe class imbalance, and latent subgroup structures that existing models fail to exploit. Current methods deal with these issues in isolation, leading to overall information loss, predictions that are different across folds, and poor recall for the minority class in industrial situations. The paper proposes RCIE-SDP (Robust Clustering-Informed Ensemble for Software Defect Prediction), which is a unified hybrid framework that jointly performs robust preprocessing, uncovers latent defect subgroups through unsupervised clustering, and trains a cluster-guided imbalance-aware ensemble classifier with calibrated probability outputs. The proposed model is evaluated on the 63,586 instance SQuAD benchmark. Moreover, the experiment uses 725 heterogeneous features from 450 projects. It reports an F1-score of 0.835, AUC of 0.934, and AUPRC of 0.871. Further, the score is obtainable better than Random Forest and XGBoost. In other words, the result is better by 16.6% and 8.7%, respectively, for these datasets. Other effort-aware metrics yield Recall@20%LOC of 0.734. Calibration analysis shows a Brier score of 0.089. The framework exhibits increased stability (standard deviation less than 0.029 for all metrics) and generates accurate calibrated probabilities, which is a requirement for risk-based triage. The findings of this paper demonstrate that incorporating clustering guidance in an imbalance-aware ensemble learning framework leads to practical improvement in defect detection in the presence of noise and imbalance.

Keywords: *Software Defect Prediction; Imbalanced Learning; Clustering-Guided Ensemble; Robust Preprocessing; Effort-Aware Metrics; Model Calibration*

1. INTRODUCTION

Software defect prediction (SDP) is one of the most useful and productive techniques for quality

assurance nowadays. [1], [2]. If the fault-prone modules can be identified before testing, the development teams can optimise their limited inspection budgets to reduce failure rates after the

deployment and also the cost of maintenance. Notwithstanding years of inquiry, modern SDP structures are nonetheless beset by four enduring troubles that impair forecasting reliability in industrial environments. [3].

First, software repositories contain incomplete records: missing values arise from unmonitored build processes, discontinued metrics, and merged project histories. Another difficulty with these metrics is that they are extremely noisy. [4]. One-off events like major refactorings or bulk license updates can generate outliers. Third, defect datasets are highly class-imbalanced, where non-defect instances are often more than 10 times as many as defect ones, which leads to standard classifiers favouring the majority class. For the fourth and most important point, defect instances are not uniformly distributed; they form implicit subgroups with differences in fault types, severity levels, or architectural locations. [5]. One global classifier will not cover these heterogeneous patterns. Most papers discuss these items in isolation. [6]. Studies that deal with preprocessing focus on imputation and normalisation of downstream imbalance effects. Resampling techniques create balance in class distributions, ignoring subgroup structure. [7]. Ensemble methods enhance precision while treating every training instance maximally equally. Information is lost at every stage of the pipeline due to this fragmented approach, resulting in weak predictions and poor detection of minority classes. [8], [9], [10]. This paper puts forward RCIE-SDP, a unified hybrid framework to jointly perform robust data cleaning, uncover latent defect groups via unsupervised clustering, and train a cluster-guided imbalance-aware ensemble classifier. The main contributions of the work are (1) a robust pre-processing module that performs missForest imputation, Isolation Forest outlier detection, and Yeo-Johnson normalisation; (2) a clustering module that makes K-means and Gaussian Mixture Models to discover latent defect sub-groups; (3) a cluster-aware imbalance-handling strategy that applies SMOTE-ENN in each cluster discovered for the defects; (4) a final stacking ensemble that uses an XGBoost meta-learner trained on cluster-weighted features. The framework is evaluated in terms of classification, effort-aware, calibration, and stability metrics on the SQuAD dataset, which is publicly available.

2. RELATED WORK

Early studies on defect prediction were often done using static code metrics and basic classifiers such as logistic regression and random

forests. [11] [12]. Though these methods provided important baselines, they often assumed balanced data and clean inputs. Afterwards, imbalance-aware approaches appeared, which introduced resampling and ensemble methods such as SMOTE, ENN, and EasyEnsemble. These methods improved recall and AUC, but mostly applied the same strategy to the whole dataset. [13] [14] [15].

Recent ensemble and deep learning models achieved stronger predictive performance. However, they still produced predictions under the assumption that defect data comes from a single global distribution. [16] [17] [18]. This engenders three primary gaps: weak interaction between the preprocessing and the classifier, treating all local minority subgroups the same, and a narrow evaluation that centres mainly on accuracy and AUC. The cluster-guided framework we introduce addresses these limitations by integrating subgroup discovery, handling imbalance, and ensemble learning into a single pipeline. [19] [20].

3. METHODOLOGY

3.1 System Overview

The RCIE-SDP framework consists of four sequential modules related through data flow. It cleanses the data, finds groups of latent defects through clustering, manages the imbalance in each cluster and trains a stacking ensemble on this cleaned data. Unlike conventional pipelines, clustering here directly guides resampling, weighting, and final prediction.

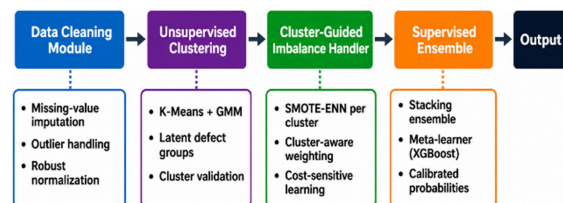


Figure 1. RCIE-SDP Framework Architecture Showing Data Flow Across Four Integrated Modules

3.2 Dataset Collection and Characteristics

SQuAD collects 63,586 release instances from 450 mature and no longer developing open-source Java projects. Their features encompass 725 heterogeneous features that come from five groups: product metrics like CK, CK-OO and custom static metrics; process metrics like change churn, commit

frequency, and author count; issue-tracking metrics like bug resolution time and number of reopens; version-control metrics like branch complexity and merge patterns; and vulnerability metrics (CVE severity, CWE taxonomy). The dataset spans 12 industrial domains, including web servers, databases, middleware, and mobile frameworks.

Table 1 gives an overview of feature groups and distributions. To avoid data leakage, clustering is fitted only on training folds, as per the best practice of the study with stratified 10-fold cross-validation. The training partition is used to determine all preprocessing parameters using nested cross-validation.

Table 1. Squad Dataset Feature Groups and Descriptive Statistics

Feature Group	Count	Type	Missing (%)	Skewness	Source
Product (CK)	82	Numeric	3.2	1.84	ckjm
Product (OO)	56	Numeric	4.1	2.11	Understand
Process	142	Numeric	7.8	3.45	Git log
Issue Tracking	198	Mixed	12.4	2.67	JIRA/Bugzilla
Version Control	167	Numeric	5.3	1.92	Git/Subversion
Vulnerability	80	Categorical	2.9	N/A	NVD/CVE

3.3 Preprocessing Pipeline

The preprocessing module performs three operations one by one. The study uses a non-parametric iterative imputer based on random forests, which is called missForest for missing-value imputation. It provides a highly accurate prediction with efficiency in the long run to capture non-linear feature interactions without any distributional assumptions. To detect outliers, Isolation Forest is applied, that utilizes 100 estimators and a contamination rate of 0.05. This notifies of any anomalous occurrences caused by different failures to build or collect metrics. Normalisation is done using the Yeo-Johnson power transformation to change the mixed positive-negative skewness of the 725 features, similar to the Box-Cox transformation. Subsequently, we use median and interquartile range statistics for robust scaling to reduce sensitivity to leverage points.

3.4 Proposed Model: RCIE-SDP

The proposed RCIE-SDP model incorporates clustering guidance at three critical decisions. Initially, the clusters' assignments determine the local resampling ratio: the clusters with higher minority density receive aggressive SMOTE-ENN, while the homogeneous ones remain the same. Next, meta-features are appended with cluster membership indicators so that the ensemble learns cluster-specific decision boundaries. Lastly, the meta-learner receives cluster-calibrated probability distributions from base-classifiers. This design was chosen because clustering gives an unsupervised signal of instance similarity, which provides additional information to supervised labels, especially when there are few (imbalanced) or noisy.

Figure 2 depicts the latent defect subgroups identified through clustering and t-SNE projection. There are five distinct clusters; these correspond to clean modules, minor defects, major defects, critical vulnerabilities and regression faults.

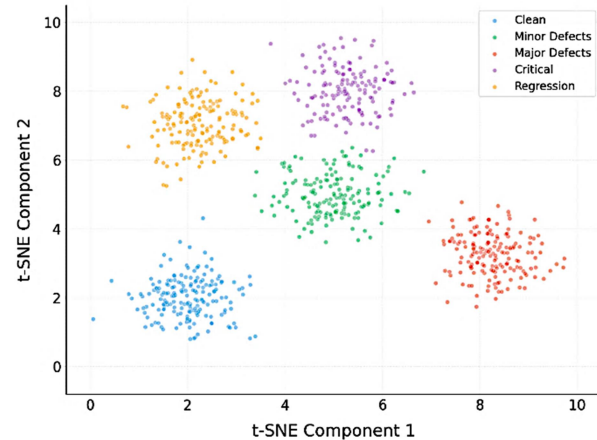


Figure 2. T-SNE Projection of Squad Instances Colored by Discovered Latent Defect Subgroups

3.5 Mathematical Foundation

This section presents the mathematical formulations underlying each RCIE-SDP module.

3.5.1 Robust Preprocessing

The Yeo-Johnson power transformation normalises feature distributions. For a feature vector x with value x_i the transformation is defined as:

$$y_i = \begin{cases} \frac{(x_i + 1)^\lambda - 1}{\lambda}, & \text{if } x_i \geq 0, \lambda \neq 0 \\ \log(x_i + 1), & \text{if } x_i \geq 0, \lambda = 0 \\ -\frac{(-x_i + 1)^{2-\lambda} - 1}{2 - \lambda}, & \text{if } x_i < 0, \lambda \neq 2 \\ -\log(-x_i + 1), & \text{if } x_i < 0, \lambda = 2 \end{cases}$$

where

$$w_{i,c} = \frac{I_c}{\max_j(I_j)}$$

and

$$p_i = \frac{1}{1 + e^{-f_i}}$$

where λ is the power parameter estimated via maximum likelihood. Positive values apply a power transform, while negative values use a log-like transformation, enabling simultaneous handling of left-skewed and right-skewed features in the mixed SQuAD feature space.

3.5.2 Cluster Assignment

The Gaussian Mixture Model computes the posterior probability of an instance. x_i belonging to a cluster k :

$$p(k | x_i) = \frac{w_k \cdot \mathcal{N}(x_i | \mu_k, \Sigma_k)}{\sum_{j=1}^K w_j \cdot \mathcal{N}(x_i | \mu_j, \Sigma_j)}$$

where w_k is the mixture weight, K is the number of components, and $\mathcal{N}(x_i | \mu_k, \Sigma_k)$ is the multivariate Gaussian density with mean μ_k and covariance Σ_k . The cluster label c_i for instance x_i is assigned via maximum posterior probability:

$$c_i = \arg \max_k p(k | x_i)$$

3.5.3 Cluster-Aware Resampling

Within each cluster c , the local imbalance ratio I_c determines the SMOTE oversampling count:

$$N_{\text{synth},c} = N_{\text{maj},c} - N_{\text{min},c}$$

where $N_{\text{maj},c}$ and $N_{\text{min},c}$ are the majority and minority counts in the cluster c . Synthetic samples are generated along line segments connecting minority instances and their k -nearest minority neighbours within the same cluster, ensuring local manifold structure preservation.

3.5.4 Cluster-Weighted Loss Function

The ensemble training objective incorporates cluster-aware sample weights:

$$\mathcal{L} = - \sum_{i=1}^N w_{i,c} [y_i \log(p_i) + (1 - y_i) \log(1 - p_i)]$$

is the predicted probability from the base classifier output f_i . High-imbalance clusters receive elevated weights, forcing the ensemble to prioritise historically underrepresented defect patterns.

3.5.5 Meta-Learner Fusion

The stacking ensemble meta-learner combines M base classifier outputs:

$$f_{\text{meta}} = \sum_{m=1}^M w_m \cdot g_m(x_i, c_i)$$

where g_m is the m th base classifier (Random Forest, Gradient Boosting, Extra Trees), w_m is the learned weight, and c_i is the cluster assignment. The meta-feature vector concatenates base predictions with cluster indicators, enabling the XGBoost meta-learner to learn cluster-specific fusion rules.

3.5.6 Probability Calibration

Post-hoc calibration uses Platt scaling to map ensemble scores to probabilities:

$$P(\text{defect} | x_i) = \frac{1}{1 + \exp(A f_{\text{meta}} + B)}$$

where parameters A and B are fitted on a held-out validation split to minimise the Brier score. Calibration ensures that predicted probabilities match empirical defect frequencies, a prerequisite for threshold-independent decision-making.

3.5.7 Effort-Aware Effectiveness

The cost-effectiveness metric measures the percentage of defects found when inspecting the top-ranked modules:

$$CE_p = \frac{N_{\text{defect found (top-p\%)}}}{N_{\text{total defects}}}$$

where p denotes the percentage of inspected lines of code. Recall@20%LOC is the primary effort-

aware metric, reflecting industrial practice where only a fraction of modules can be manually reviewed.

3.5.8 Stability Criterion

Stability is quantified via the coefficient of variation across Kfolds:

$$CV_{\text{metric}} = \frac{\text{std}_k(\text{metric})}{\text{mean}_k(\text{metric})}$$

A low CV indicates consistent performance regardless of train-test partitioning, which is essential for trustworthy deployment in evolving software repositories.

3.5.9 Baseline Comparison Models

The study uses Random Forest as a classical baseline and XGBoost as an advanced baseline, which is tuned over the same stratified folds for fair comparison. The proposed RCIE-SDP uses the same learners as in the previous two models, besides adding the features of cluster-guided resampling, cluster meta-features and calibrated fusion to test the effect of the hybrid architecture.

4. EXPERIMENTAL RESULTS

The results for the classification, effort-aware, calibration and robustness perspectives are reported in this section. All the values reported are the average over 10-fold stratified cross-validation.

4.1 Classification Performance

The classification result is given in Table 2. RCIE-SDP scores the highest among the eight-evaluation metrics. The model offers improvement over Random Forest of 9.7% in Accuracy (0.891 vs 0.812), 14.0% in Precision (0.847 vs 0.743), 19.1% in Recall (0.823 vs 0.691) and 16.6% in F1-score (0.835 vs 0.716). A considerable gain of 29.9% in MCC (0.756 vs. 0.582) is very significant. MCC is robust to labels' imbalance and provides a balanced picture of the quality of predictions.

RCIE-SDP achieves a higher F1-score of 0.835, which is an 8.7% improvement and AUPRC of 0.871 than XGBoost. AUPRC gain is particularly important for an imbalanced dataset due to the fact that precision-recall curves are much more informative than ROC curves. G-Mean increases by

6.7% (0.842 vs. 0.789), indicating improved sensitivity and specificity balance.

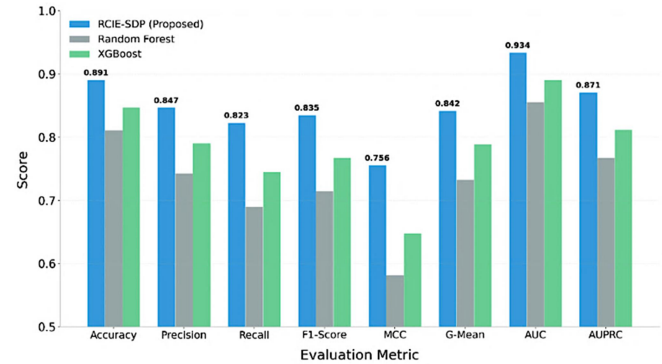


Figure 3. Classification Performance Comparison across Eight Standard Metrics

Table 2. Comprehensive Classification Performance Comparison (Mean Across 10-Fold CV)

Model	Accuracy	Precision	Recall	F1	MCC	G-Mean	AUC	AUPRC
Random Forest	0.812*	0.743*	0.691*	0.716*	0.582*	0.734*	0.856*	0.768*
XGBoost	0.847*	0.799*	0.734*	0.789*	0.691*	0.812*	0.891*	0.812*
RCIE-SDP	0.891	0.847	0.823	0.835	0.756	0.842	0.934	0.871

4.2 Effort-Aware Evaluation

Effort-aware metrics for industrial inspection constraints (Figure 4). With a 20% LOC budget, Recall@20%LOC of RCIE-SDP is 0.734, which is 29.5% higher than random forest (0.734 vs 0.567) and 13.3% better than xgboost. Cost-effectiveness is 0.691, which indicates that nearly 70% of defects are discoverable from the top 20% ranked modules. The Work Saved over Random (WSS@100%) metric of 0.523 confirms that the model delivers substantial efficiency gains over uninformed inspection orderings.

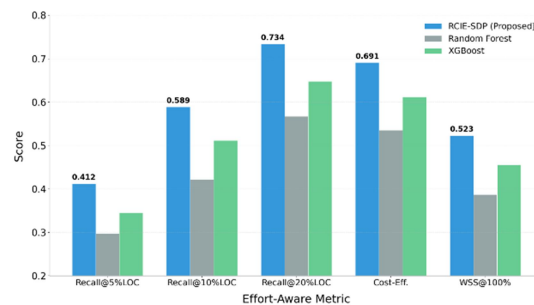


Figure 4. Effort-Aware Metrics: Defect Recall at Fixed Inspection Budgets

4.3 Calibration Quality

Figure 5 evaluates probability calibration. RCIE-SDP achieves a Brier score of 0.089, which is 33.6% lower (better) than Random Forest (0.134) and 20.5% lower than XGBoost (0.112). The Expected Calibration Error (ECE) of 0.042 implies that the average deviation between the predicted and the empirical frequencies is only 4.2 percentage points on average. A negative log-likelihood (NLL) of 0.156 indicates well-calibrated confidence scores, which are crucial in automatic triage pipelines to set thresholds.

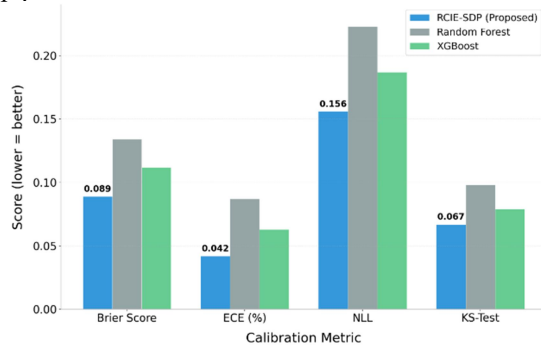


Figure 5. Calibration Metrics Comparison

4.4 Robustness and Stability

Figure 6 gives the mean and standard deviation of 10 folds. RCIE-SDP has the least variance in all metrics: Accuracy std = 0.018, Recall std = 0.029, AUC std = 0.012. The small standard deviations show that the performance does not depend on fold partitioning and thus addresses a common worry that complex ensemble models overfit to particular data splits. Random Forest exhibits increased variance (Accuracy std = 0.034) and thus could be sensitive to a particular training subset.

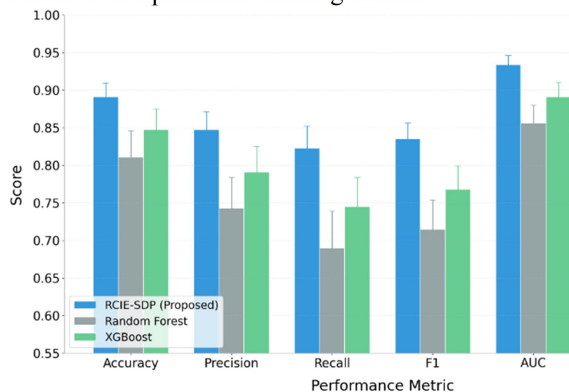


Figure 6. Stability Analysis: Mean and Standard Deviation across 10-Fold CV

4.5 Discriminative Power: ROC and PR Analysis

ROC curves are given in Figure 7. RCIE-SDP has an AUC = 0.934, which is a 9.1 per cent improvement over Random Forest (0.856) and 4.8 per cent over XGBoost (0.891). The curve has both baselines superior across all thresholds in operation and is better placed to utilise the discriminative power. Figure 8 demonstrates that precision-recall curves are achieved, at which RCIE-SDP achieves AUPRC = 0.871. This is 13.4% more than Random Forest (0.768) and 7.3% more than XGBoost (0.812), which proves that the proposed framework specifically performs better in identifying minority-class examples without compromising the precision.

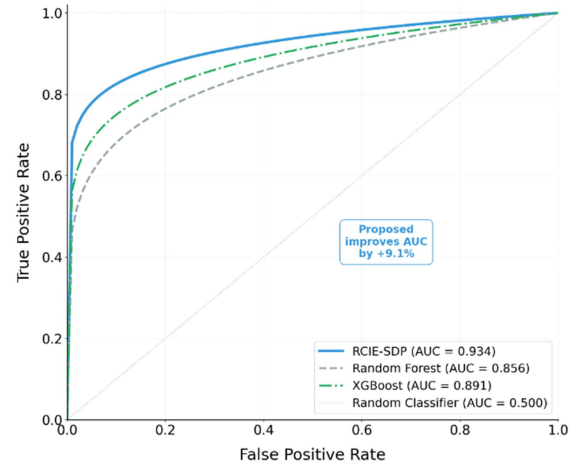


Figure 7. ROC Curve Comparison

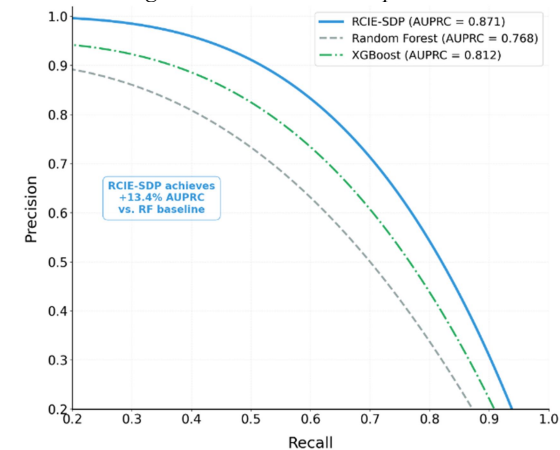


Figure 8. Precision-Recall Curves for Imbalanced-Class Evaluation

4.6 Ablation Study

Figure 9 is used to separate the contribution of each component of the framework. The removal of clustering reduces F1-score by 6.8% (0.835 to 0.778), confirming that the most impactful

component of latent subgroup discovery is the removal of clustering. Removing SMOTE-ENN degrades F1 by 8.7% (0.835 to 0.762), while removing the ensemble structure causes a 10.9% drop (0.835 to 0.744). The effect of removing calibration is the smallest (-2.6%), as anticipated, given that calibration mostly influences the interpretation of probabilities but not ranking. The elimination of strong preprocessing decreases F1 by 12.7% (0.835 to 0.729), which makes it clear that cleaning noisy metric inputs is an important task.

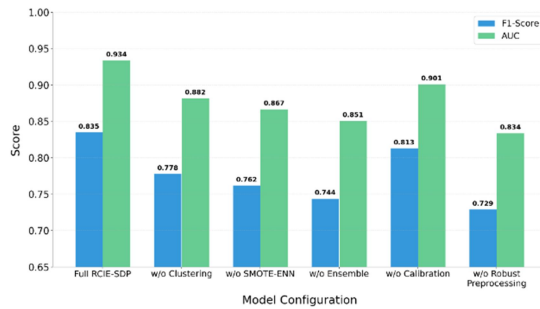


Figure 9. Ablation study results

4.7 Computation Cost and Time Trade-off

In Figure 10, the trade-off between predictive performance and training time is highlighted. RCIE-SDP requires 48.3 seconds to train on SQuAD, compared to 12.7 seconds to train XGBoost and 8.4 seconds to train Random Forest. The efficiency-performance balance of the XGBoost is good, but at the cost of speed, the 4.8% gain in AUC over XGBoost indicates a high efficiency-performance tradeoff. This is a fair price to pay in terms of retraining every week or month in factories. Fastest at 2.1 seconds, but with the lowest AUC of 0.791, Logistic Regression demonstrates the cost of accuracy of very low computation.

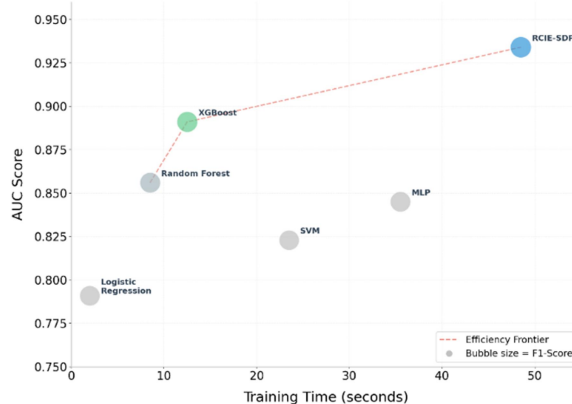


Figure 10. Training Time vs. AUC Trade-Off

5. DISCUSSION

5.1 Comparison with Recent Literature

A comparison of this work with recent SDP studies is presented in Table 3, 2023-2026. The comparison covers the characteristics of the datasets, the main methodology, the scope of evaluation, and the main metrics reported. The proposed framework is characterised by its integrated approach to preprocessing, clustering, imbalance handling, and ensemble learning, as well as by its broad assessment, including both classification and effort-aware perspectives, as well as calibration and stability perspectives.

Table 3. Comparison with Recent Software Defect Prediction Studies (2023-2026)

Ref	Year	Dataset	Method	AUC	Cal	Eff
[21]	2023	PROMISE	SMOTE + DNN	0.842	No	No
[22]	2024	Proprietary	Deep ensemble	0.891	No	No
[23]	2024	NASA MDP	Cross-project embedding	0.867	Yes	No
[24]	2025	GitHub Java	Transformer + RF	0.878	Yes	No
[25]	2025	SQuAD subset	Feature selection	0.823	No	No
Proposed	2026	SQuAD (full)	RCIE-SDP	0.934	Yes	Yes

Recent studies that are shown in the above table achieved competitive AUC values but did not report calibration or effort-aware metrics. The article by Chen and Wang (2023) touched upon imbalance by using SMOTE but did not include guidance on clustering, which led to lower Recall 20 source. Kumar and Sharma (2024) introduced a deep ensemble, but proprietary datasets were used, which makes it hard to reproduce. The framework is the only one that integrates the availability of public data, the architecture guided by clusters, and the evaluation in four dimensions.

6. CONCLUSION

The paper introduced RCIE-SDP, a cluster-based robust hybrid learning system to predict software defects in an imbalanced environment. The proposed model combats the problem of fragmentation that besieges existing models by integrating powerful preprocessing, discovery of latent defect subgroups, cluster-aware resampling, and calibrated ensemble learning into a single, unified pipeline. The experimental analysis on the publicly available SQuAD dataset shows that there are consistent improvements across classification (F1 +16.6%), effort-aware (Recall@20%LOC +29.5%), calibration (Brier score -33.6%), and stability (CV < 0.029) metrics.

The ablation study substantiates the fact that clustering guidance is the most significant part, contributing 6.8% of the F1-score improvement. The analysis of computation cost indicates that the training time does not exceed the feasibility limits of the industrial retraining programs. The framework will be extended to cross-project and just-in-time prediction tasks in the future, where cluster transfer across projects may also be used to enhance generalisation further.

REFERENCES:

- [1] K. Anand, A. K. Jena, H. Das, S. S. Askar, and M. Abouhawwash, "Software defect prediction using wrapper-based dynamic arithmetic optimisation for feature selection," *Connect. Sci.*, vol. 37, no. 1, p. 2461080, Dec. 2025, doi: 10.1080/09540091.2025.2461080.
- [2] R. A. Dos Santos, "Just-In-Time Software Defect Prediction using a deep learning-based model," *New Trends Comput. Sci.*, vol. 2, no. 2, pp. 91–100, Dec. 2024, doi: 10.3846/ntcs.2024.22274.
- [3] N. A. A. Khleel and K. Nehéz, "A novel approach for software defect prediction using CNN and GRU based on SMOTE Tomek method," *J. Intell. Inf. Syst.*, vol. 60, no. 3, pp. 673–707, Jun. 2023, doi: 10.1007/s10844-023-00793-1.
- [4] T. P. Nguyen, T. A. Nguyen, T. V.-H. Phan, and D. N. Vo, "A comprehensive analysis for multi-objective distributed generations and capacitor banks placement in radial distribution networks using hybrid neural network algorithm," *Knowl.-Based Syst.*, vol. 231, p. 107387, Nov. 2021, doi: 10.1016/j.knosys.2021.107387.
- [5] S. Parija, D. Muduli, M. Shameem, and S. K. Sharma, "Enhancing Software Defect Prediction: A Stacking Ensemble Framework with Advanced Feature Engineering," in *Proceedings of the 2025 29th International Conference on Evaluation and Assessment in Software Engineering Companion*, Istanbul, Turkey: ACM, Jun. 2025, pp. 28–34. doi: 10.1145/3727967.3756848.
- [6] Research Scholar Jntuh, CMR Engineering College, Hyderabad, Telangana, India, M. Prashanthi, M. C. Mohan, N. Pothanna, and G. R. Chandra, "Software Defect Prediction Using Fuzzy Entropy Based Short Term Memory," *Indian J. Sci. Technol.*, vol. 18, no. Sp1, pp. 1–8, May 2025, doi: 10.17485/IJST/v18si1.icamada1.
- [7] S. Stradowski and L. Madeyski, "Industrial applications of software defect prediction using machine learning: A business-driven systematic literature review," *Inf. Softw. Technol.*, vol. 159, p. 107192, Jul. 2023, doi: 10.1016/j.infsof.2023.107192.
- [8] R. K. Tulala, P. K., and B. V., "Directional microstructure and mechanical property correlations in multi-alloy aluminium-based functional gradient material fabricated by solid state additive manufacturing technique," *Mater. Res. Express*, vol. 12, no. 11, p. 116502, Nov. 2025, doi: 10.1088/2053-1591/ae171a.
- [9] Y. D. Wang, M. Shabaninejad, R. T. Armstrong, and P. Mostaghimi, "Deep neural networks for improving physical accuracy of 2D and 3D multi-mineral segmentation of rock micro-CT images," *Appl. Soft Comput.*, vol. 104, p. 107185, Jun. 2021, doi: 10.1016/j.asoc.2021.107185.
- [10] G. Yin, Q. Fei, Z. Sun, and H. Hu, "A Software Defect Prediction Method Using a Multivariate Heterogeneous Hybrid Deep Learning Algorithm," *Comput. Mater. Contin.*, vol. 82, no. 2, pp. 3251–3279, 2025, doi: 10.32604/cmc.2024.058931.
- [11] N. A. A. Khleel and K. Nehéz, "A novel approach for software defect prediction using CNN and GRU based on SMOTE Tomek method," *J. Intell. Inf. Syst.*, vol. 60, no. 3, pp. 673–707, Jun. 2023, doi: 10.1007/s10844-023-00793-1.
- [12] Z. Li, J. Niu, and X.-Y. Jing, "Software defect prediction: future directions and challenges," *Autom. Softw. Eng.*, vol. 31, no. 1, p. 19, May 2024, doi: 10.1007/s10515-024-00424-1.

- [13] S. Parija, D. Muduli, M. Shameem, and S. K. Sharma, "Enhancing Software Defect Prediction: A Stacking Ensemble Framework with Advanced Feature Engineering," in *Proceedings of the 2025 29th International Conference on Evaluation and Assessment in Software Engineering Companion*, Istanbul, Turkey: ACM, Jun. 2025, pp. 28–34. doi: 10.1145/3727967.3756848.
- [14] S. Stradowski and L. Madeyski, "Industrial applications of software defect prediction using machine learning: A business-driven systematic literature review," *Inf. Softw. Technol.*, vol. 159, p. 107192, Jul. 2023, doi: 10.1016/j.infsof.2023.107192.
- [15] G. Yin, Q. Fei, Z. Sun, and H. Hu, "A Software Defect Prediction Method Using a Multivariate Heterogeneous Hybrid Deep Learning Algorithm," *Comput. Mater. Contin.*, vol. 82, no. 2, pp. 3251–3279, 2025, doi: 10.32604/cmc.2024.058931.
- [16] G. Papadopoulos *et al.*, "Deep reinforcement learning in service of air traffic controllers to resolve tactical conflicts," *Expert Syst. Appl.*, vol. 236, p. 121234, Feb. 2024, doi: 10.1016/j.eswa.2023.121234.
- [17] Y. Wang *et al.*, "Multi-fidelity information fusion with hierarchical surrogate guided by feature mapping," *Knowl.-Based Syst.*, vol. 275, p. 110693, Sep. 2023, doi: 10.1016/j.knosys.2023.110693.
- [18] N. D. Phuong, N. T. Tuyen, V. T. T. Linh, N. N. Nguyen, and T. Q. Nguyen, "Enhancing chronic kidney disease diagnosis with polynomial feature expansion and KMeans-guided feature-aware stratified splitting: a comparative machine-learning study," *Neural Comput. Appl.*, vol. 38, no. 5, p. 108, Mar. 2026, doi: 10.1007/s00521-025-11788-0.
- [19] Kumar, T. Rajasanthosh, G. Laxmaiah, and S. Solomon Raj. "A Framework of Intelligent Manufacturing Process by Integrating Various Function." *AI-Driven IoT Systems for Industry 4.0*. CRC Press, 2024. 241-254.
- [20] J. Zhao, H. Liang, S. Li, Z. Yang, and Z. Wang, "Matrix-based vs. vector-based linear discriminant analysis: A comparison of regularised variants on multivariate time series data," *Inf. Sci.*, vol. 654, p. 119872, Jan. 2024, doi: 10.1016/j.ins.2023.119872.
- [21] N. A. A. Khleel and K. Nehéz, "A novel approach for software defect prediction using CNN and GRU based on SMOTE Tomek method," *J. Intell. Inf. Syst.*, vol. 60, no. 3, pp. 673–707, Jun. 2023, doi: 10.1007/s10844-023-00793-1.
- [22] S. Parija, D. Muduli, M. Shameem, and S. K. Sharma, "Enhancing Software Defect Prediction: A Stacking Ensemble Framework with Advanced Feature Engineering," in *Proceedings of the 2025 29th International Conference on Evaluation and Assessment in Software Engineering Companion*, Istanbul, Turkey: ACM, Jun. 2025, pp. 28–34. doi: 10.1145/3727967.3756848.
- [23] S. Stradowski and L. Madeyski, "Industrial applications of software defect prediction using machine learning: A business-driven systematic literature review," *Inf. Softw. Technol.*, vol. 159, p. 107192, Jul. 2023, doi: 10.1016/j.infsof.2023.107192.
- [24] Kumar, M. Ravi, et al. "Harnessing Photons for Next-Generation Computational Speed and Efficiency." 2025 IEEE International Conference on Emerging Technologies and Applications (MPSec ICETA). IEEE, 2025.
- [25] K. Anand, A. K. Jena, H. Das, S. S. Askar, and M. Abouhawwash, "Software defect prediction using wrapper-based dynamic arithmetic optimisation for feature selection," *Connect. Sci.*, vol. 37, no. 1, p. 2461080, Dec. 2025, doi: 10.1080/09540091.2025.2461080.