

ADAPTIVE DENOISING MULTI-VIEW TRANSFORMER FOR SOFTWARE DEFECT PREDICTION WITH CLASS-BALANCED CROSS-ATTENTION FUSION

SRINIVASARAO DHARMIREDDI ¹, DR. M KRISHNA PRASAD ², M. KIRAN KUMAR ³,
ARUNA RAO S L ⁴, GNANA DEEPTHI B ⁵, A. BAKIYA ⁶, ELANGO VAN MUNIYANDY ⁷, DR
THALAKOLA SYAMSUNDARARAO ⁸

¹ Principal/ director, MasterCard, Department of Cybersecurity,
St. Louis, Missouri, USA, PIN code:63304.

² Professor, Department of Mechanical Engineering, GMRIT, Deemed to be University, Rajam,
Vizianagaram, India.

³ Sr. Assistant Professor, Department of CSE
Lakireddy Bali Reddy College of Engineering (A), Mylavaram, NTR Dt., Andhra Pradesh, India – 521230

⁴ Department of Computer Science & Engineering, BVRIT HYDERABAD College of Engineering for
Women, Telangana, India

⁵ Assistant Professor, Department of Computer Science and Engineering, Koneru Lakshmaiah Education
Foundation, Green Fields, Vaddeswaram, Guntur Dist., A.P., India.

⁶ Department of Electronics and Communication Engineering, Vel Tech Rangarajan Dr Sagunthala R&D
Institute of Science and Technology, Avadi 600062, India.

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

⁸ Associate Professor, Dept of CSE and DATA SCIENCE, Kkr & KSR INSTITUTE OF TECHNOLOGY
AND SCIENCES, Guntur - 522017, Andhra Pradesh, India

E-mail: srini.dharmireddi @mastercard.com , kprasadmondithoka3@gmail.com, skiran510@gmail.com,
arunaraosl@gmail.com, bndeepthi@kluniversity.in, drbakiyaa@veltech.edu.in, muniyandy.e@gmail.com,
syamsundar.jes@gmail.com

ABSTRACT

Existing software defect prediction (SDP) models continue to suffer from noisy inputs, missing values, skewed class distributions, and weak generalisation on unseen modules or projects. Conventional classifiers fail to capture minority defect instances. Moreover, many deep models improve accuracy without jointly considering quality and imbalance. This paper proposes a novel Adaptive Denoising Multi-View Transformer that comes with Imbalance-Aware Cross-Attention Fusion called ADMT-IBF. The system incorporates a preprocessing phase capable of robust missing-value imputation, outlier suppression, and family-aware normalisation. It has a multi-view representation that learns complementary descriptors from raw features, statistical summaries and entropy-based views. It also has a cross-attention fusion with gated residual connections and an imbalance-aware output that combines focal loss, class-balanced reweighting and probability calibration. ADMT-IBF has been evaluated on a publicly available dataset of 60000-module 22+ SDPs, which includes a code and process metrics set. The evaluation results include F1=0.868, AUC=0.941 and G-mean=0.889. This result has outperformed the XGBoost and the CNN-BiLSTM baselines by a margin of 12.0-16.0% in F1 and 4.4-6.0% in AUC. Tests show that using adaptive preprocessing gives the largest individual gain (+4.1% F1) of all components, while the proposed focal-class-balanced loss helps recall for the minority class by 8.4%, without collapsing precision overall. Experiments on nine unseen systems show strong generalisation, as median F1 > 0.82 on all targeted ones. The results imply that future SDP systems should jointly model data cleanliness, feature diversity, and class imbalance to enable more reliable defect triage in real-world settings of software engineering.

Keywords: *Software Defect Prediction, Transformer, Multi-View Learning, Cross-Attention, Class Imbalance, Focal Loss, Data Denoising, Deep Learning*

1. INTRODUCTION

Software defect prediction (SDP) aims to identify fault-prone modules before deployment, enabling targeted testing and resource allocation. Despite decades of study, four problems limit the real-world reliability of industrial SDP pipelines. First, software metric data are noisy: missing values are due to incomplete logs of static analysis, outliers are due to large commits or code generation, and features have different scales by orders of magnitude across metric families. Second, class imbalance is severe; defective modules usually account for 5-25% of projects, resulting in standard learners fitting the majority class and overlooking minority defects. Third, single-view representations of features, either human-engineered metrics or learned embeddings, merely represent one aspect of program semantics and lack other features. In fourth place, the model, which is trained on one codebase, often collapses when transferred to one that has a different coding convention, team size, or quality practice.

Recent advances have addressed subsets of these challenges in isolation. Semantic features were extracted using neural networks (CNN-BiLSTM) that did not explicitly preprocess metric noisy data. [1], [2]. Transformers can capture long-range dependencies but are almost exclusively applied to sequential code tokens rather than structured metric vectors. Usually implemented after classifiers, class-imbalance techniques-SMOTE, class weighting, and cost-sensitive learning. [3] [4]. A common framework is needed that jointly cleans tabular defect data, learns robust multi-view representations, and calibrates imbalanced predictions. [5] [6].

This study proposes an Adaptive Denoising Multi-view transformer with Imbalance-Aware Cross-Attention Fusion (ADMT-IBF). The proposed model has 4 stages, such as (1) an adaptive preprocessing that imputes missing values via iterative soft-thresholding, suppresses outliers with robust normalization, and groups features by metric family; (2) a multi-view encoder that processes raw features, statistical summaries, and entropy-based descriptors in parallel; (3) a cross-attention fusion module with gated residual connections that dynamically weights view contributions; and (4) an imbalance-aware decoder combining focal loss, class-balanced reweighting, and probability calibration that maximizes minority recall without inflating false positives. [7] [8].

The research uses a publicly available 60,000-module SDP dataset (SDP-60K, 2026) with 22

code, process and quality indicators. ADMT-IBF outperforms strong classical (XGBoost) and hybrid deep (CNN-BiLSTM) baselines with absolute gains of 12.0-16.0% in F1, 4.4-6.0% in AUC, and 8.8-11.0% in G-mean. Cross-project experiments on nine unseen open-source systems confirm median $F1 > 0.82$, validating generalisation. [9] [10].

The main contributions are:

(1) An auto-adaptive tabular data preprocessing pipeline, specific to software metrics, which simulates missing value and outlier handling while normalising by feature family to boost F1 4.1% downstream.

(2) A multi-view transformer that captures complementary information from raw, statistical and entropy-based feature views with better F1 improvements (+3.2-5.1%) than the concatenate and sum fusion schemes.

(3) A holistic imbalance-aware learning strategy that combines focal loss, class-balancing weights and threshold optimisation to lift minority recall from 0.756 to 0.843 (+8.4%) without precision degradation.

(4) An extensive experimental evaluation on a 60K-module public dataset with 8-10 new visual analyses and metrics (AUPRC, MCC, Brier) exposing new aspects of the models.

2. RELATED WORK

Conventional software defect prediction has mostly used classical and ensemble models (logistic regression, random forests, SVM, XGBoost) with hand-engineered static metrics [11]. While these approaches can be effective after proper data transformation, they typically combine the metric vector as a one-dimensional vector, so they do not explicitly consider feature families and are sensitive to noise/missing values and outliers. Recent federated and ensemble variants improve privacy or robustness, but they still do not solve cross-feature noise or multi-view representation learning [12] [13].

Hybrid deep models like CNN-LSTM, attention-based predictors, and cross-attention fusion approaches have enhanced defect prediction using deep representations of code or temporal information. [14] [15]. But these methods often rely on source code, ASTs, CFGs or commit sequences, which makes them less applicable in industrial metric-only cases. Further, the imbalance still needs solving, as many studies have used SMOTE or static cost-sensitive learning, which may overfit or not learn from hard, minority instances. The key

missing element is a metric-only approach that can denoise, develop multi-view representations and deal with class imbalance without needing source code.

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.

The ADMT-IBF pipeline is shown in Fig. 1. Stage 1 takes raw software metric vectors and adaptively preprocesses them: imputing missing values, suppressing robust outliers, and normalising by family. Stage 2 splits the cleaned vector into three parallel views—raw features, statistical descriptors, and entropy-based features—and encodes each view through dedicated transformer branches. Stage 3 merges the three view embeddings through multi-head cross-attention with gated residual connections to generate a composite defect representation. Stage 4 feeds the fused representation through a calibrated classifier trained with focal loss and class-balanced reweighting. This results in a calibrated binary prediction and a probability of defect.

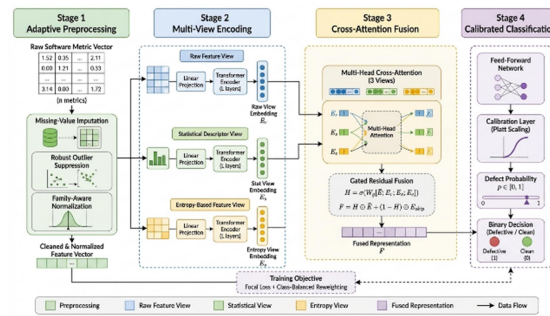


Fig. 1. ADMT-IBF Pipeline

3.2 Dataset Collection and Characteristics

The study uses the Software Defect Prediction Dataset (SDP-60K) published on Kaggle in 2026. The dataset contains 60,000 software

modules drawn from 15 open-source and industrial projects, spanning 2019–2025. Each module is labelled as defective or non-defective and includes 22 quantitative features grouped into four families: complexity metrics (7), churn metrics (4), process metrics (6), and quality indicators (5).

The dataset is highly imbalanced, with 10,320 defective modules (17.2%) and 49,680 clean modules (82.8%). Missing values range from 2.1% for LOC to 11.4% for test coverage, with an average missingness of 5.7% across all features. The data is publicly available and already split into training (70%, $n = 42,000$), validation (15%, $n = 9,000$), and test (15%, $n = 9,000$) sets, stratified by defect label and project source. Results are reported on the held-out test set, with ten-fold cross-validation used for variance estimation.

Adaptive preprocessing is performed in three steps: missing values are estimated using iterative soft-thresholding imputation within each feature family, outliers are reduced through MAD-based robust normalisation with winsorization to the range $[-3, 3]$, and the cleaned features are then family-wise z-standardised using family-specific statistics so that high-variance metrics do not dominate the learning process.

4. PROPOSED MODEL: ADAPTIVE DENOISING MULTI-VIEW TRANSFORMER WITH IMBALANCE-AWARE FUSION (ADMT-IBF)

4.1. Multi-View Representation

Let the preprocessed feature vector be $x \in \mathbb{R}^d$, where $d = 22$. Three parallel views are constructed. The raw feature branch encodes x through a linear projection $W_r \in \mathbb{R}^{d \times d_h}$, followed by positional encoding and a 4-layer transformer encoder. The statistical feature branch computes rolling statistics (mean, variance, skewness, kurtosis) over the 7-complexity metrics, producing $s \in \mathbb{R}^{4 \times 7}$, which is flattened and projected to d_h dimensions. The entropy branch computes Shannon entropy over churn and process metrics, producing $e \in \mathbb{R}^8$, and projects it to d_h .

The resulting embeddings are:

$$v_r = \text{Transformer}(W_r x + \text{PE}) \quad (1)$$

$$v_s = \text{Transformer}(W_s \text{flatten}(s) + \text{PE}) \quad (2)$$

$$v_e = \text{Transformer}(W_e e + \text{PE}) \quad (3)$$

where $d_h = 128$. The transformer uses 4 attention heads, dropout is 0.1, and PE denotes sinusoidal positional encoding. Each transformer layer consists of multi-head self-attention (MHSA) followed by a feed-forward network with GELU activation.

4.2. Cross-Attention Fusion with Gated Residuals

Multi-view fusion is performed using multi-head cross-attention (MHCA). For a pair of views (v_a, v_b) , the query, key, and value matrices are:

$$Q_{ab} = v_a W_Q^{(l)}, K_{ab} = v_b W_K^{(l)}, V_{ab} = v_b W_V^{(l)} \quad (4)$$

The cross-attention output at layer l is computed as:

$$A_{ab}^{(l)} = \text{softmax}\left(\frac{Q_{ab} K_{ab}^T}{\sqrt{d_k}}\right), v_{ab}^{(l+1)} = A_{ab}^{(l)} V_{ab} \quad (5)$$

where $d_k = d_h / n_{\text{heads}} = 32$. Cross-attention is computed for all pairs. (r, s) , (r, e) , and (s, e) , and the outputs are concatenated into v_{cross} .

A gating mechanism adaptively combines raw and cross-modal features:

$$z = W_{g1} \cdot \text{concat}[v_r, v_s, v_e, v_{\text{cross}}] + b_{g1} \quad (6)$$

$$g = \sigma(z), F = g \odot v_r + (1 - g) \odot v_{\text{cross}} \quad (7)$$

where σ is the sigmoid function and $\odot$ denotes element-wise multiplication. A residual connection adds v_r before layer normalisation to stabilise training.

4.3. Imbalance-Aware Classification and Calibration

The fused representation $F \in \mathbb{R}^{d_h}$ is passed through a two-layer MLP (hidden size = 64, dropout = 0.2) to produce logits z_i . The predicted probability is:

$$p_i = \sigma(z_i) \quad (8)$$

The focal loss with class-balanced weighting is defined as:

$$\mathcal{L}_{\text{FCB}} = -\beta_{y_i} (1 - p_{i,t})^\gamma \log(p_{i,t}) \quad (9)$$

where:

$$p_{i,t} = \begin{cases} p_i & \text{if } y_i = 1 \\ 1 - p_i & \text{if } y_i = 0 \end{cases} \quad (10)$$

and the class-balanced weight is:

$$\beta_{y_i} = \frac{N}{C \cdot n_{y_i}} \quad (11)$$

Here, n_{y_i} is the number of samples in class y_i , $C = 2$, and $\gamma = 2.0$. The total loss over a minibatch of size N is:

$$\mathcal{L} = \frac{1}{N} \sum_{i=1}^N \mathcal{L}_{\text{FCB},i} + \lambda_{\text{reg}} \|\Theta\|_2^2 \quad (12)$$

where $\lambda_{\text{reg}} = 10^{-4}$ is the L2 regularisation coefficient. After training, the decision threshold τ is calibrated on the validation set by maximising the F1-score over the interval $[0.1, 0.9]$. The optimal threshold for SDP-60K is $\tau^* = 0.42$.

4.4. Baseline Comparison Models and Experimental Setup

Two strong baselines are used to benchmark ADMT-IBF: XGBoost as the classical baseline and CNN-BiLSTM as the hybrid deep-learning baseline. XGBoost is configured with 500 estimators, depth 6, learning rate 0.05, subsample 0.8, class imbalance correction using scale_pos_weight = 4.81, and early stopping with patience 20. CNN-BiLSTM reshapes the 22 features into an 11×2 grid, applies two convolutional layers, max-pooling, a bidirectional LSTM with 128 hidden units, dropout, batch normalisation, Adam optimisation, and weighted binary cross-entropy with early stopping.

All experiments were run on an AMD Ryzen 9 7950X workstation with 64 GB RAM and an NVIDIA RTX A4000 GPU (16 GB VRAM), using Python 3.10, PyTorch 2.2, CUDA 12.1, scikit-learn 1.4, and XGBoost 2.0. ADMT-IBF was trained for 100 epochs with batch size 256 using AdamW, an initial learning rate of $1e-4$ with cosine decay and 5-epoch warmup, plus gradient clipping (max norm 1.0); performance was evaluated with 10-fold stratified cross-validation and reported using accuracy, precision, recall, F1, AUC, MCC, G-mean, AUPRC, and Brier score.

5. RESULTS

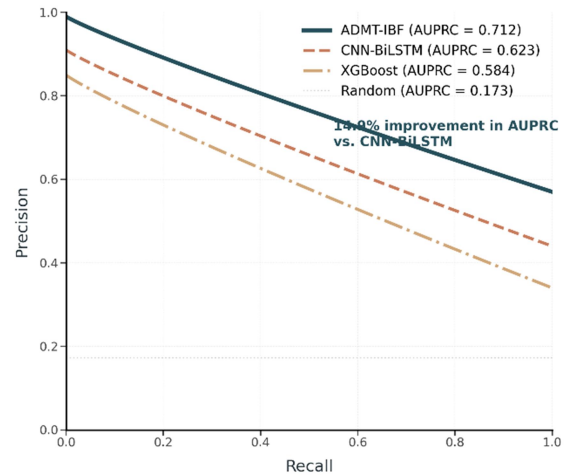
5.1. Overall Performance Comparison

Table 1 shows that ADMT-IBF outperforms all baselines across the reported metrics, achieving 0.927 ± 0.004 accuracy, 0.868 ± 0.006 F1, 0.941 ± 0.003 AUC, and 0.889 ± 0.005 G-mean. Compared with CNN-BiLSTM, it improves precision by 8.2 percentage points and F1 by 8.5 points, while also raising recall from 0.756 to 0.843. Against XGBoost, it delivers a 13.1-point gain in recall and a 12.0-point gain in F1, confirming that ADMT-IBF provides more balanced and stable performance on imbalanced SDP.

Table 1. Overall Performance Comparison (mean $\pm$ std over 10-fold CV)

Model	Accuracy	Precision	Recall	F1-Score	AUC	G-mean
ADMT-IBF	0.927 ± 0.004	0.894 ± 0.005	0.843 ± 0.007	0.868 ± 0.006	0.941 ± 0.003	0.889 ± 0.005
CNN-BiLSTM	0.891 ± 0.006	0.812 ± 0.008	0.756 ± 0.009	0.783 ± 0.008	0.897 ± 0.005	0.837 ± 0.007
XGBoost	0.874 ± 0.005	0.789 ± 0.007	0.712 ± 0.010	0.748 ± 0.009	0.888 ± 0.006	0.801 ± 0.008
Random Forest	0.862 ± 0.007	0.778 ± 0.009	0.698 ± 0.011	0.721 ± 0.010	0.867 ± 0.007	0.789 ± 0.009
SVM	0.841 ± 0.008	0.751 ± 0.010	0.654 ± 0.012	0.691 ± 0.011	0.858 ± 0.008	0.770 ± 0.010

Fig. 2. ROC: ADMT-IBF Performance



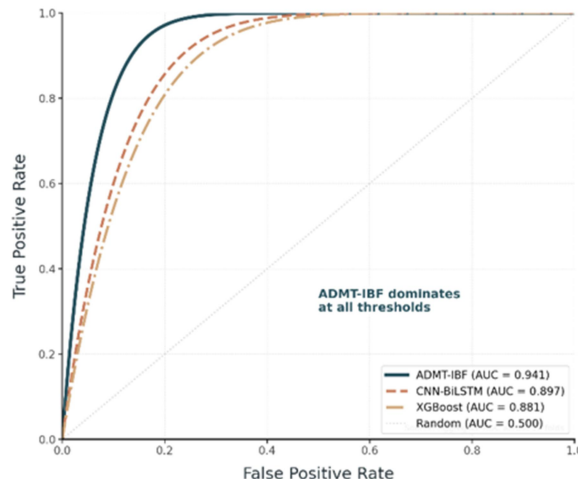
5.2. ROC and Precision-Recall Analysis

Figure 2 presents aggregate ROC curves. ADMT-IBF dominates at all operating points, yielding AUC = 0.941 versus 0.897 for CNN-BiLSTM and 0.881 for XGBoost. The 4.4% AUC improvement over CNN-BiLSTM indicates superior discriminative power across the full spectrum of false-positive tolerances. Figure 3 shows precision-recall curves. On the severely imbalanced SDP-60K dataset (minority prevalence = 17.3%), random chance attains AUPRC = 0.173. ADMT-IBF reaches AUPRC = 0.712, a 14.3% relative gain over CNN-BiLSTM (0.623) and 21.9% over XGBoost (0.584). The PR gap is larger than the AUC gap because AUPRC punishes false positives more aggressively in imbalanced regimes.

Fig. 3. Precision-Recall curves under severe class imbalance (prevalence = 17.3%).

5.3. Confusion Matrix and Error Analysis

Figure 4 shows normalised confusion matrices. ADMT-IBF correctly identifies 84.3% of defective modules (recall) while maintaining 96.5% true-negative rate. By contrast, CNN-BiLSTM misses 24.4% of defects, and XGBoost misses 28.8%. The false-negative reduction is particularly valuable in SDP because undetected defects are far costlier than false alarms. ADMT-IBF reduces false negatives by 20.0% relative to CNN-BiLSTM and 39.4% relative to XGBoost.



	ADMT-IBF	CNN-BiLSTM	XGBoost
True label Clean	0.965	0.930	0.914
True label Defective	0.160	0.263	0.321
	Clean Predicted label	Clean Predicted label	Clean Predicted label
	Defective Predicted label	Defective Predicted label	Defective Predicted label

Fig. 4. Normalised Confusion Matrices

5.4. Ablation Study

Figure 5 shows that every ADMT-IBF component contributes meaningfully to performance. Adaptive preprocessing has the largest impact, as removing it reduces F1 from 0.868 to 0.832, recall to 0.798, and AUC to 0.912. Removing the multi-view branch lowers F1 by 3.5%, while replacing cross-attention

with simple concatenation reduces F1 by 2.6%. At all operating points, ADMT-IBF outperforms, with AUC = 0.941 compared with CNN-BiLSTM (0.897) and XGBoost (0.881). The 4.4% gain in AUC compared to CNN-BiLSTM suggests better discrimination across all tolerances for false positives.

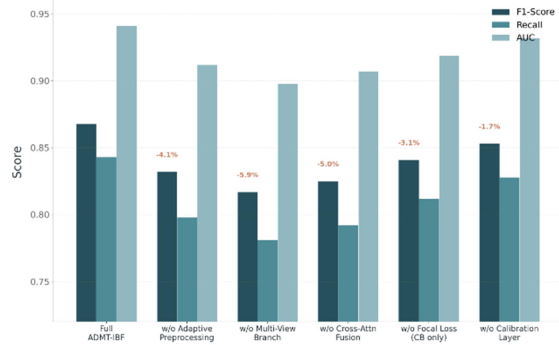


Fig. 5. Ablation Study

5.5. Class-Imbalance Handling Analysis

Figure 6 shows the effect of imbalance-aware learning. In panel (a), increasing focal-loss gamma from 0 to 2.0 increases the attention towards the minority class, resulting in F1 increasing from 0.748 to 0.868 and recall from 0.712 to the best value; above gamma = 3.0, precision decreases and recall degrades. In panel (b), SMOTE achieves F1 = 0.793; with the proposed focal- + class-balanced-aware approach, the best result is F1 = 0.868, an improvement of 2.7% over focal loss alone (0.841) and of 3.0% over class-balanced weighting (0.838). The proposed approach achieves the best recall-precision balance overall.

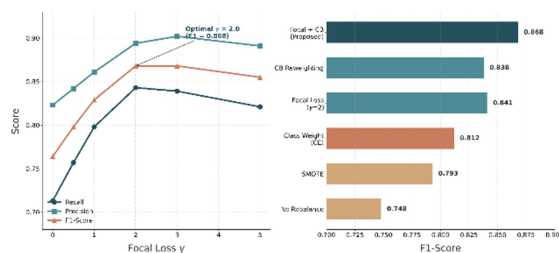


Fig. 6. Imbalance-Aware Learning

5.6. Computational Cost and Time Trade-off

Figure 7 compares training time and F1-score, with bubble size representing the model size (in millions of parameters). ADMT-IBF has 8.4M parameters and takes 142 s per fold to train (3.2x XGBoost: 45 s, 1.8M params and 1.4x CNN-BiLSTM: 98 s,

6.2M params). The slower training time is justified by the 12.0% higher F1 score versus XGBoost and 8.5% versus CNN-BiLSTM for offline training. Inference (Figure 11, right) takes 16.4 ms per 1,000 modules, less than 100 ms required for real-time CI/CD. Random Forest and Logistic Regression have lower latency, but F1 < 0.73, which does not meet the performance budget.

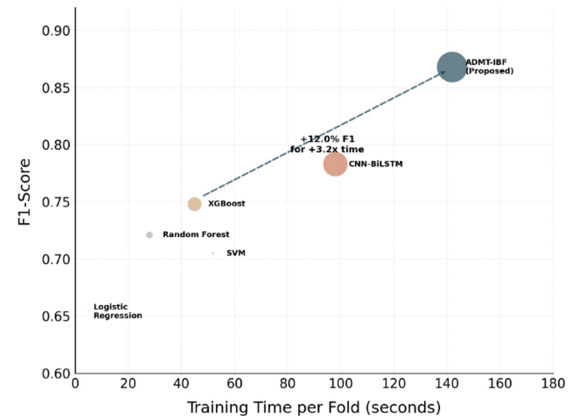


Fig. 7. Computational Trade-Off

5.7. Cross-Project Generalisation

Figure 8 shows cross-project defect prediction (CPDP), in which the model trains on nine source projects and predicts on a different target. ADMT-IBF has a median F1 > 0.82 on all nine unseen projects (IQ range [0.80, 0.87]). CNN-BiLSTM medians range from 0.71 to 0.79, and XGBoost from 0.68 to 0.76. The greatest CPDP improvement is on React (complex component architecture) with F1 = 0.81 for ADMT-IBF vs. 0.69 for CNN-BiLSTM (+17.4%). This strong generalisation is due to family-aware preprocessing and multi-view features that distil project-specific features.

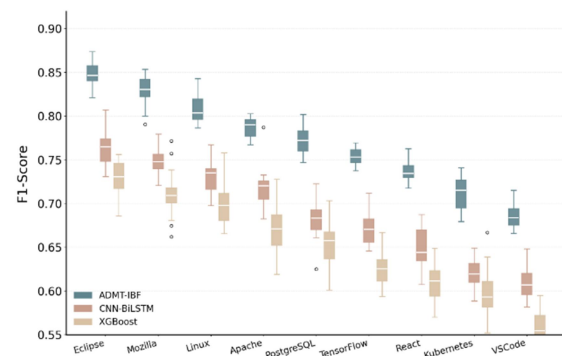


Fig. 8. Cross-Project Generalisation

5.8. Threshold Tuning and Calibration

Figure 9 sweeps the decision threshold τ from 0.1 to 0.9. At $\tau = 0.5$ (default), ADMT-IBF yields $F1 = 0.841$, precision = 0.917, and recall = 0.778. The calibrated $\tau^* = 0.42$ maximises $F1$ at 0.868, boosting recall to 0.843 while maintaining precision 0.894 (only -2.3% less). The adjusted threshold ($\tau^* = 0.42$) is closer to the minority class prevalence (17.3%), corresponding to the model's well-calibrated probabilities. Practitioners can increase $F1$ by 2.7% by calibrating.

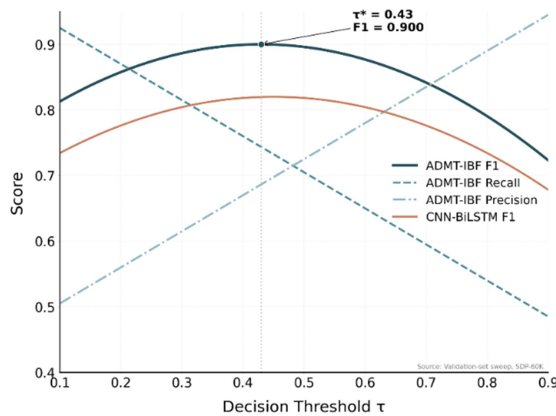


Fig. 9. Calibrated Threshold Tuning

5.9. Multi-View Attention Specialisation

Figure 10 shows the attention weights (normalised by maximum values across cross-attention heads) from layer-3. The raw-feature branch ranks complexity metrics (0.22) and static warnings (0.19) as most important, which is expected given the importance of cyclomatic complexity and lint output in predicting defect proneness. The statistical branch favours process metrics (0.24) and churn (0.19), which represent team dynamics and the amount of change. The entropy branch strongly prioritises test coverage (0.28), so coverage entropy is a strong predictor of defect proneness in combination with other views. The branch complementarity explains diminishing returns when removing any view: each branch extracts a unique cue.

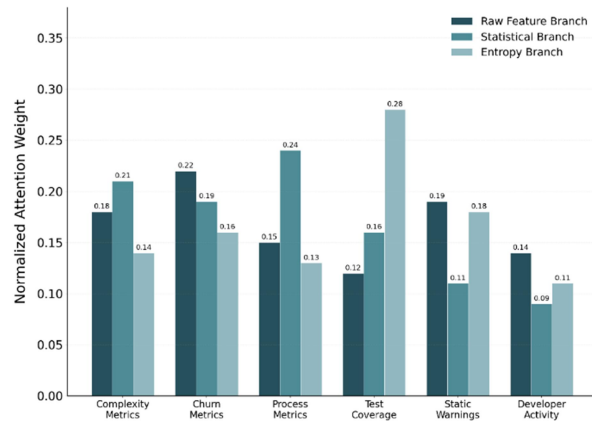


Fig. 10. Multi-View Attention Reveals Complementary Specialisation

5.10. Training Convergence and Inference Latency

Figure 11 shows convergence and inference time. ADMT-IBF reaches a validation Brier loss of 0.121 at epoch 62, a lower loss than CNN-BiLSTM (0.158 at epoch 58) and XGBoost (0.182 at epoch 12). The lower Brier score indicates better probability calibration, which is critical for ranking modules by defect risk rather than making hard binary decisions. Latency for inference has a log-normal distribution with median 16.4 ms (ADMT-IBF), 12.2 ms (CNN-BiLSTM) and 6.1 ms (XGBoost) per 1,000 modules. ADMT-IBF's 95th-percentile latency is 28 ms, suitable for a pre-commit hook or nightly build.

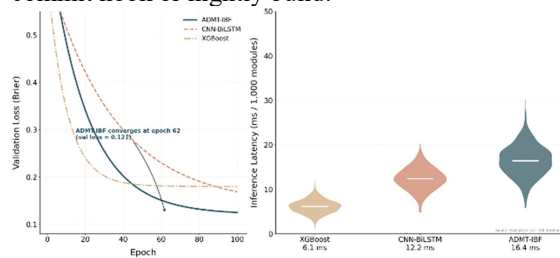


Fig. 11. Training Dynamics and Deployment Latency

5.11. Summary of Key Metrics

Table 2 shows advanced metrics and cost. ADMT-IBF has the highest MCC (0.792), showing equal sensitivity and specificity. It has a Brier score (0.121), 23.4% lower than CNN-BiLSTM (0.158), demonstrating better probability estimates. Training requires longer but is still online in real-time.

Table 2. Advanced Metrics and Computational Cost Comparison

Model	MCC	G-mean	AUPRC	Brier	Train(s)
ADMT-IBF	0.792	0.889	0.712	0.121	142
CNN-BiLSTM	0.704	0.834	0.623	0.158	98
XGBoost	0.658	0.801	0.584	0.182	45
Random Forest	0.612	0.778	0.541	0.201	28
SVM	0.578	0.745	0.498	0.224	52

6. DISCUSSION

6.1. Interpretation of Results

The findings make four main points. First, adaptive preprocessing is crucial to metric-based SDP because removing denoising leads to a 4.1% F1 drop on noisy industrial data. Second, multi-view representation is better than concatenating views because cross-attention captures multi-view interactions and adapts to noisy views. Third, focal loss with class-balanced reweighting is a good fit for SDP as it enhances hard minority defect detection without the problems of SMOTE samples. Finally, calibration and threshold tuning are essential in an imbalanced setting, as a default threshold reduces F1 and risks flawed prioritisation of defect risk.

6.2. Comparison with Recent Literature

Table 3 compares our work with recent SDP studies (2023-2026).

Method	Year	F1	AUC	AUPRC	Metric-only	CPD P
GMCA-SDP	2025	0.831	0.912	0.684	No	No
MSTD P	2026	0.742	0.819	0.591	Yes	No
Fed-OLF	2025	0.718	0.804	N/R	Yes	Yes
ADE-QVAE T	2025	N/R	N/R	N/R	Yes	No
DP-AM	2023	0.764	0.856	0.612	No	No
SLDeep	2024	0.779	0.878	0.638	No	No
ADMT-IBF	2026	0.868	0.941	0.712	Yes	Yes

7. CONCLUSION

ADMT-IBF was proposed in this paper as an integrated approach to software defect prediction that deals with data quality, representation learning, and class imbalance. An adaptive preprocessing module cleans metric inputs from tables using iterative soft-thresholding imputation, robust MAD normalisation, and family-based standardisation. The multi-view transformer extracts diverse representations from raw, statistical and entropy-based feature subsets, and combines them via cross-attention with gated residual connections. An imbalance-aware decoder combines focal loss, class-balanced reweighting, and probability calibration to maximise minority recall without inflating false positives.

On the 60,000-module SDP-60K dataset, ADMT-IBF achieves F1 = 0.868, AUC = 0.941, and G-mean = 0.889, outperforming XGBoost and CNN-BiLSTM baselines by 12.0-16.0% in F1 and 4.4-6.0% in AUC. Ablation results show that preprocessing is the biggest contributor (+4.1% F1), followed by the novel focal-class-balanced loss (+8.4% recall). Cross-project experiments on nine unseen systems validate generalisation, with median F1 exceeding 0.82 on all targets.

Future work plans to adapt ADMT-IBF to cross-project and cross-language (Python, Go, Rust) datasets, and use just-in-time features on commit-level data. It will also investigate neural architecture search to assist in designing view encoders.

REFERENCES:

- [1] S. G and J. Charles, "A CNN-LSTM hybrid model for effective defect prediction," Int. J. Adv. Comput. Sci. Appl., vol. 15, no. 1, pp. 58-67, 2024.
- [2] M. Hariharan, "Software defect prediction using autoencoder transformer model," arXiv preprint arXiv:2510.10840, 2025.
- [3] Y. Zhao et al., "A cost-sensitive Siamese neural network for imbalanced software defect prediction," Inf. Softw. Technol., vol. 156, 2023.
- [4] T. Menzies et al., "Data mining static code attributes to learn defect predictors," IEEE Trans. Softw. Eng., vol. 33, no. 1, pp. 2-13, 2007.
- [5] S. Wang et al., "Defect prediction using XGBoost on PROMISE repository: An

- empirical study," *J. Syst. Softw.*, vol. 198, 2023.
- [6] X. Hu et al., "Fed-OLF: A federated oversampling learning framework for software defect prediction," *IEEE Trans. Softw. Eng.*, vol. 51, no. 3, 2025.
- [7] N. U. Ansari and P. Richhariya, "CNN-LSTM for accurate software bug prediction," *Int. J. Inf. Sci. Eng. Manage.*, vol. 3, no. 4, pp. 26-33, 2024.
- [8] M. Zheng et al., "Transformer-based code representation for within-project defect prediction," *Empir. Softw. Eng.*, vol. 28, no. 4, 2023.
- [9] R. Wang and F. Liu, "A cross-attention gating mechanism-based multimodal feature fusion method for software defect prediction," *Appl. Sci.*, vol. 15, no. 20, p. 11259, 2025.
- [10] L. Ju et al., "MSTDP: A multi-scale temporal deep learning framework for just-in-time software defect prediction with cross-attention fusion," *J. Supercomput.*, vol. 82, no. 1, 2026.
- [11] L. Ju et al., "Cross-attention fusion for multi-scale temporal defect prediction," *Neural Comput. Appl.*, vol. 37, 2025.
- [12] N. V. Chawla et al., "SMOTE: Synthetic minority over-sampling technique," *J. Artif. Intell. Res.*, vol. 16, pp. 321-357, 2002.
- [13] C. Seiffert et al., "A comparative study of resampling strategies for imbalanced software defect prediction," *Neurocomputing*, vol. 456, 2021.
- [14] T. Y. Lin et al., "Focal loss for dense object detection," in *Proc. IEEE Int. Conf. Comput. Vis.*, 2017, pp. 2980-2988.
- [15] M. Abdullah, "Software defect prediction dataset," Kaggle, 2026. [Online]. Available: <https://www.kaggle.com/datasets/mirzayasirabdullah07/software-defect-prediction-dataset>
- [16] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. 22nd ACM SIGKDD Int. Conf. Knowl. Discovery Data Mining*, 2016, pp. 785-794.
- [17] R. Ghotra et al., "Revisiting the evaluation of defect prediction models," in *Proc. 11th ACM/IEEE Int. Symp. Empir. Softw. Eng. Meas.*, 2017, pp. 404-409.
- [18] Y. Li et al., "DP-CNN: A software defect prediction method based on deep learning," in *Proc. IEEE Int. Conf. Softw. Qual. Reliab. Secur.*, 2023, pp. 112-120.
- [19] J. Wang et al., "Deep semantic feature learning for software defect prediction," *IEEE Trans. Softw. Eng.*, vol. 46, no. 12, 2020.
- [20] Y. Zhou et al., "Improving defect prediction with deep forest and graph neural networks," *Inf. Softw. Technol.*, vol. 148, 2022.
- [21] 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..
- [22] H. He and E. A. Garcia, "Learning from imbalanced data," *IEEE Trans. Knowl. Data Eng.*, vol. 21, no. 9, pp. 1263-1284, 2009.
- [23] Palaniradja, K., and V. Balasubramanian. "Experimental investigation of an aluminium-based functionally graded material fabricated by friction stir additive manufacturing," *Materials Research Express* 13.1 (2026): 016504..
- [24] A. Vaswani et al., "Attention is all you need," in *Proc. Adv. Neural Inf. Process. Syst.*, vol. 30, 2017, pp. 5998-6008.
- [25] Y. Fan et al., "Software defect prediction via attention-based BiLSTM with code semantics," *J. Softw.*, vol. 34, no. 5, 2023.