

Soft computing techniques for defect prediction and software quality improvement

Tamanna

Assistant Professor, Department of Computer Science, Chhotu Ram Arya College, Sonipat, Haryana, India

DOI: <https://doi.org/10.66856/ijraet.2026.12.3.12025>

Abstract

Software defects increase development cost, delay release schedules, and reduce user satisfaction. Conventional defect prediction methods often rely on rigid assumptions, linear relationships, and manually selected thresholds. Soft computing provides an alternative by modeling uncertainty, imprecision, and nonlinear relationships in software engineering data. This paper reviews the application of fuzzy logic, artificial neural networks, evolutionary computation, support vector machines, swarm intelligence, and hybrid methods to software defect prediction and quality improvement. The study analyzes the role of software metrics, data preprocessing, feature selection, model training, and performance evaluation. A conceptual framework is proposed in which repository data and process metrics are transformed into defect-risk estimates and quality-improvement recommendations. The analysis indicates that hybrid soft computing models can improve prediction capability when compared with individual techniques, particularly in the presence of class imbalance and noisy project data. However, issues related to interpretability, data quality, cross-project generalization, and model maintenance remain significant. The paper concludes that soft computing should be integrated with software quality assurance processes rather than treated as an isolated prediction mechanism.

Keywords: Soft computing, software defect prediction, software quality, fuzzy logic, neural networks, evolutionary computation, machine learning, software metrics.

Introduction

Software quality is a major concern in modern information systems because defects can produce financial loss, security vulnerabilities, service interruptions, and safety risks. The increasing size and complexity of software systems make exhaustive manual inspection impractical. Consequently, software organizations require automated methods for identifying defect-prone modules before deployment.

Software defect prediction uses historical project information to estimate whether a software component is likely to contain one or more defects. Typical prediction variables include lines of code, cyclomatic complexity, code churn, coupling, inheritance depth, developer activity, historical fault count, and testing information. Although these variables are useful, their relationships with defects are rarely simple. A large module may be defective because of its size, but a small module may also be defective because of high coupling or frequent modification.

Traditional statistical methods generally assume a predefined relationship between independent variables and defect occurrence. Such assumptions may not accurately represent software development data. Software repositories also contain missing values, inconsistent measurements, duplicated records, noisy labels, and severe class imbalance. These characteristics motivate the use of soft computing techniques, which are designed to tolerate uncertainty and approximate complex relationships.

This paper has the following objectives:

- To examine major soft computing techniques used for software defect prediction.
- To describe how these techniques support software quality improvement.
- To compare their advantages and limitations.
- To propose a general framework for integrating prediction models with quality assurance activities.
- To identify research challenges and future directions.

Background

a. Software Defects and Quality

A software defect is an error, fault, or failure-producing condition in a software artifact. Defects may originate in requirements, design, implementation, configuration, or maintenance activities. In practical systems, defects are not uniformly distributed. A relatively small proportion of modules may contain a large proportion of observed defects. Software quality includes several attributes, such as:

- Functional correctness.
- Reliability.
- Maintainability.
- Security.
- Performance efficiency.
- Usability.
- Testability.

Defect prediction primarily supports reliability and maintainability, but its effect can extend to other quality attributes. For example, identifying highly coupled modules may support both defect reduction and maintainability improvement.

b. Software Metrics

Defect prediction models depend on measurable characteristics of software projects. Common metric categories include:

1. **Product metrics:** lines of code, complexity, coupling, cohesion, inheritance, and code duplication.
2. **Process metrics:** number of revisions, code churn, bug-fix frequency, and historical defect density.
3. **Developer metrics:** number of contributors, ownership changes, and developer experience.
4. **Repository metrics:** issue activity, commit frequency, pull-request history, and review outcomes.

5. **Testing metrics:** statement coverage, branch coverage, failed test cases, and regression-test history.

Let a software module be represented by a feature vector

$$\mathbf{x}_i = [x_{i1}, x_{i2}, \dots, x_{im}] \quad (1)$$

Where x_{ij} denotes the value of metric for module i . The target variable y_i may be binary:

$$y_i = \begin{cases} 1, & \text{if module } i \text{ is defective,} \\ 0, & \text{otherwise.} \end{cases}$$

The prediction model estimates the probability

$$\hat{p}_i = P(y_i = 1 \mid \mathbf{x}_i) \quad (2)$$

A module is classified as defect-prone when $\hat{p}_i$ exceeds a selected decision threshold.

c. Limitations of Conventional Prediction Methods

Conventional methods face several difficulties:

- Defect data are often highly imbalanced.
- Metric distributions differ between projects.
- Historical labels may be incomplete or incorrect.
- The relationship between metrics and defects is nonlinear.
- A model trained on one project may perform poorly on another.
- Binary predictions do not always communicate the degree of risk.
- Managers may require explanations rather than only numerical outputs.

Soft computing addresses some of these limitations through approximate reasoning, adaptive learning, and population-based optimization.

Soft Computing Techniques

a. Fuzzy Logic

Fuzzy logic represents variables using degrees of membership rather than strict binary categories. For example, code complexity can be classified as low, medium, or high with overlapping membership functions.

For a metric, a triangular membership function can be written as

$$\mu(z) = \begin{cases} 0, & z \leq a, \\ \frac{z-a}{b-a}, & a < z \leq b, \\ \frac{c-z}{c-b}, & b < z < c, \\ 0, & z \geq c. \end{cases} \quad (3)$$

A fuzzy defect prediction system may use rules such as:

- If complexity is high and code churn is high, then defect risk is very high.
- If coupling is medium and historical defects are low, then defect risk is medium.
- If testing coverage is high and code churn is low, then defect risk is low.

The output is typically a continuous risk value. Fuzzy systems are valuable because their rules can be interpreted by developers and quality managers. Their main limitation

is that rule design and membership-function selection may require expert knowledge.

b. Artificial Neural Networks

Artificial neural networks learn nonlinear mappings between software metrics and defect labels. A feed-forward network with one hidden layer can be represented as

$$\hat{y} = g\left(\sum_{k=1}^h w_k f\left(\sum_{j=1}^m v_{kj} x_j + b_k\right) + c\right) \quad (4)$$

Where f and g are activation functions, v_{kj} and w_k are learned weights, b_k and c are bias terms, and h is the number of hidden neurons.

Neural networks can model interactions among complexity, churn, coupling, and historical defects without requiring an explicit mathematical relationship. Deep neural networks may also process source-code representations, commit messages, and textual issue descriptions.

However, neural networks require sufficient training data and careful regularization. Their predictions may be difficult to interpret, which can reduce trust in safety-critical or regulated environments.

c. Support Vector Machines

Support vector machines classify modules by constructing a separating boundary with a maximum margin. For linearly separable data, the decision function is

$$f(\mathbf{x}) = \mathbf{w}^T \mathbf{x} + b \quad (5)$$

For nonlinear data, a kernel function can map the input into a higher-dimensional feature space. The radial basis function kernel is commonly expressed as

$$K(\mathbf{x}_i, \mathbf{x}_j) = \exp\left(-\gamma \|\mathbf{x}_i - \mathbf{x}_j\|^2\right) \quad (6)$$

Where γ controls the influence of individual training examples.

Support vector machines often perform well with small or medium-sized datasets and high-dimensional feature spaces. Their disadvantages include sensitivity to parameter selection, difficulty in interpreting the resulting decision boundary, and the need for appropriate scaling of input metrics.

d. Evolutionary Computation

Evolutionary computation uses mechanisms inspired by biological evolution to search for effective solutions. Genetic algorithms can optimize:

- Feature subsets.
- Model parameters.
- Fuzzy membership functions.
- Classification thresholds.
- Ensemble-model weights.

A genetic algorithm begins with a population of candidate solutions. Each solution is evaluated using a fitness function. Selection, crossover, and mutation generate successive populations.

A multi-objective fitness function may combine predictive and operational requirements:

$$F = \alpha(1 - \text{Recall}) + \beta(1 - \text{Precision}) + \gamma C \quad (7)$$

Where C represents model complexity and α , β and γ are nonnegative weights.

Evolutionary methods are useful when the search space is discontinuous or contains many competing objectives. Their computational cost and dependence on parameter settings are important limitations.

e. Swarm Intelligence

Swarm intelligence algorithms, including particle swarm optimization and ant colony optimization, solve optimization problems through interactions among simple agents. In defect prediction, these algorithms can select relevant metrics and optimize classifiers.

For particle, the velocity and position updates in particle swarm optimization may be represented as

$$v_i(t+1) = \omega v_i(t) + c_1 r_1 (p_i - x_i(t)) + c_2 r_2 (g - x_i(t)) \quad (8)$$

$$x_i(t+1) = x_i(t) + v_i(t+1) \quad (9)$$

Where p_i is the best position found by particle i , g is the global best position, and, ω , c_1 , c_2 , r_1 and r_2 are algorithm parameters.

Swarm-based feature selection can reduce irrelevant and redundant metrics. Nevertheless, optimization can become expensive for large datasets, and the resulting model may still lack interpretability.

f. Rough Sets

Rough set theory supports feature reduction and rule extraction without requiring prior assumptions about probability distributions. It separates objects into lower and upper approximations based on the available attributes.

In software defect prediction, rough sets can identify indispensable metrics and produce decision rules from historical data. Their effectiveness decreases when datasets contain substantial noise or continuous values that have not been discretized appropriately.

g. Hybrid Soft Computing Models

Hybrid models combine complementary techniques. Examples include:

- Fuzzy neural networks for adaptive and interpretable prediction.
- Genetic algorithms with support vector machines for feature and parameter optimization.
- Particle swarm optimization with neural networks.
- Fuzzy systems with ensemble classifiers.
- Neural networks combined with attention mechanisms for source-code analysis.

The main motivation for hybridization is to combine learning capability, optimization strength, and interpretability. A hybrid model should be justified by a measurable improvement rather than by algorithmic complexity alone.

Proposed Framework

A general framework for soft-computing-based defect prediction is shown conceptually below.

a. Data Collection

Data can be collected from version-control systems, issue trackers, continuous-integration servers, testing frameworks,

and static-analysis tools. Each record should be associated with a software module and a time interval.

Temporal ordering is essential. Features available before a release should be used to predict defects discovered after that release. Using future information during training creates data leakage and produces overly optimistic results.

b. Data Preprocessing

Preprocessing should include:

- Removal of duplicated records.
- Treatment of missing values.
- Detection of inconsistent labels.
- Normalization of numerical variables.
- Encoding of categorical variables.
- Removal of data leakage.
- Separation of training and testing periods.

For imbalanced data, the training set may be adjusted using class weighting, controlled oversampling, under sampling, or synthetic data generation. Any resampling operation should be applied only to the training data.

c. Feature Selection

Feature selection reduces dimensionality and can improve generalization. A selected feature subset may be defined as

$$S \subseteq \{1, 2, \dots, m\} \quad (10)$$

Where only metrics indexed by S are supplied to the prediction model.

Feature selection can be performed using correlation analysis, information gain, rough sets, genetic algorithms, particle swarm optimization, or recursive elimination. The selected metrics should also be evaluated for practical availability and interpretability.

d. Prediction and Risk Ranking

The model produces a risk score for each module. Instead of assigning equal testing priority to all modules, quality teams can rank modules according to predicted risk:

$$R_i = \hat{p}_i I_i \quad (11)$$

Where R_i is the priority score, $\hat{p}_i$ is the predicted defect probability, and I_i is the estimated impact of failure.

This approach distinguishes a frequently used but low-impact module from a high-risk module that supports a critical business function.

e. Quality Improvement Actions

Prediction results should be connected to specific engineering actions:

- Prioritize high-risk modules for code review.
- Increase unit and integration testing.
- Perform targeted refactoring.
- Review recent changes with high code churn.
- Assign experienced developers to critical components.
- Strengthen regression testing.
- Monitor defect trends after corrective actions.

The model should be periodically retrained because software architecture, development practices, and team composition change over time.

Evaluation Methodology

a. Experimental Design

A reliable evaluation should use historical project data divided into training, validation, and testing periods. A time-aware split is generally more realistic than randomly mixing records from all releases.

For cross-project evaluation, the model is trained on one or more projects and tested on a separate project. This assesses generalization across different development environments.

b. Performance Measures

Accuracy is often misleading when defective modules are rare. More informative measures include precision, recall, specificity, score, area under the receiver operating characteristic curve, and area under the precision-recall curve.

The principal measures are defined as

$$\text{Precision} = \frac{TP}{TP+FP} \quad (12)$$

$$\text{Recall} = \frac{TP}{TP+FN} \quad (13)$$

$$F_1 = \frac{2 \cdot \text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}} \quad (14)$$

Where **TP**, **FP**, and **FN** denote true positives, false positives, and false negatives.

For defect prediction, recall is important because failing to identify a defective module may allow a fault to reach production. Precision is also important because excessive false alarms waste testing and review resources.

c. Cost-Sensitive Evaluation

Software organizations may assign different costs to false positives and false negatives. A cost-sensitive objective can be expressed as

$$\text{Cost} = C_{FN}FN + C_{FP}FP \quad (15)$$

Where **C_{FN}** is the cost of missing a defective module and **C_{FP}** is the cost of investigating a nondefective module.

The decision threshold should therefore be selected according to project risk, testing capacity, and business impact rather than using a universal value.

d. Statistical Validation

Results should be reported over repeated runs or multiple project releases. Confidence intervals and statistical significance tests can support reliable comparisons. Researchers should also report computational cost, training time, feature count, and model complexity.

Comparative Analysis

Technique	Principal Strength	Main Limitation	Suitable Use
Fuzzy logic	Interpretable rules and uncertainty handling	Rule design may be subjective	Risk explanation and expert-guided prediction
Neural networks	Strong nonlinear modeling capability	Limited interpretability and data requirements	Large datasets and complex metric relationships
Support vector machines	Effective in high-dimensional spaces	Parameter sensitivity	Small or medium-sized metric datasets
Genetic algorithms	Flexible optimization and feature selection	High computational cost	Feature and parameter optimization
Swarm intelligence	Efficient population-based search	Possible premature convergence	Metric selection and classifier tuning
Rough sets	Rule extraction and attribute reduction	Sensitivity to noise and discretization	Compact rule-based models
Hybrid models	Combines complementary advantages	Increased complexity	High-performance predictive systems

No single technique is universally optimal. The most appropriate method depends on dataset size, metric type, class imbalance, required interpretability, and available computational resources.

Applications to Software Quality Improvement

a. Risk-Based Testing

Defect prediction enables test teams to allocate effort according to estimated risk. High-risk modules can receive deeper path testing, additional integration tests, and broader regression coverage.

b. Refactoring Prioritization

A module with high defect probability and high complexity may be selected for refactoring before less risky components. This supports a data-driven maintenance strategy.

c. Continuous Integration

Prediction models can be integrated into continuous-integration pipelines. Each commit may trigger the extraction of updated metrics and the generation of risk

indicators. A build may receive additional inspection when the predicted risk exceeds a project-specific threshold.

d. Release Readiness Assessment

Aggregated risk scores can support release decisions. For example, managers may monitor the proportion of high-risk modules, unresolved defects, code churn, and test coverage before approving deployment.

e. Developer Decision Support

Explainable fuzzy rules, feature importance scores, and local explanations can help developers understand why a module was classified as risky. Explanations should be presented alongside predictions to reduce the possibility of inappropriate or blind reliance on the model.

Challenges

a. Data Quality

Historical datasets may contain incomplete defect reports, inconsistent module identifiers, and inaccurate timestamps. Model quality cannot exceed the quality of the data used for training.

b. Class Imbalance

In many projects, nondefective modules substantially outnumber defective modules. A classifier may achieve high accuracy while failing to detect important defects. Evaluation should therefore emphasize recall, precision-recall analysis, and cost-sensitive measures.

c. Cross-Project Generalization

Metrics may have different meanings across programming languages, architectures, organizations, and development processes. Transfer learning, domain adaptation, normalization, and project-specific calibration can improve generalization.

d. Interpretability

High-performing models may be difficult to explain. This is a major concern when predictions affect release decisions, staffing, or safety-critical testing. Interpretable models and post hoc explanation techniques should be evaluated for faithfulness, not only convenience.

e. Concept Drift

Software systems evolve. A model trained on earlier releases may become inaccurate because of architectural changes, new developers, different testing processes, or changing defect patterns. Drift detection and periodic retraining are necessary.

f. Privacy and Security

Repository data may contain proprietary source code, personal developer information, or sensitive issue descriptions. Data collection and model deployment should follow organizational security and privacy policies.

Future Research Directions

Future research should investigate the following areas:

1. **Explainable soft computing:** Develop models that provide accurate predictions and actionable explanations.
2. **Online learning:** Update models continuously as new commits, tests, and defect reports become available.
3. **Multimodal prediction:** Combine static code metrics, source-code representations, commit messages, developer activity, and test results.
4. **Transfer learning:** Adapt models across projects while reducing dependence on large labeled datasets.
5. **Cost-aware prediction:** Optimize testing decisions according to business impact and engineering effort.
6. **Human-in-the-loop systems:** Allow developers and quality engineers to validate, correct, and refine model outputs.
7. **Security defect prediction:** Extend conventional defect prediction to vulnerabilities and security weaknesses.
8. **Green software engineering:** Include energy consumption and resource efficiency in quality-risk analysis.
9. **Reproducible benchmarks:** Establish standardized datasets, temporal evaluation protocols, and common performance measures.
10. **Hybrid optimization:** Combine interpretable fuzzy reasoning with adaptive neural and evolutionary learning.

Conclusion

Soft computing techniques provide effective tools for software defect prediction because they can represent uncertainty, learn nonlinear relationships, and optimize complex decision processes. Fuzzy logic contributes interpretability, neural networks provide adaptive nonlinear modeling, support vector machines offer strong classification performance, and evolutionary and swarm methods support feature and parameter optimization. Hybrid models can combine these strengths, but their increased complexity must be justified through rigorous evaluation.

A practical defect prediction system should include reliable data collection, leakage-free preprocessing, feature selection, cost-sensitive evaluation, risk ranking, and integration with testing and maintenance activities. Prediction alone does not improve software quality; improvement occurs when prediction results lead to appropriate engineering actions. Therefore, soft computing should be implemented as part of a continuous software quality assurance process that combines automated analysis with professional judgment.

References

1. Zadeh LA. Fuzzy sets. *Information and Control*,1965:8(3):338-353.
2. Vapnik V. *The Nature of Statistical Learning Theory*, 2nd ed. New York, NY, USA: Springer, 2000.
3. Bishop CM. *Pattern Recognition and Machine Learning*. New York, NY, USA: Springer, 2006.
4. Khoshgoftaar TM, Latifur EB, Seliya N. A brief survey of software quality prediction models. *Proc. IEEE Int. Conf. Tools with Artificial Intelligence*, 2004, 1-8.
5. Menzies T, Greenwald J, Frank A. Data mining static code attributes to learn defect predictors. *IEEE Transactions on Software Engineering*,2007:33(1):2-13.
6. Kamei Y, *et al.* A large-scale empirical study of just-in-time quality assurance. *IEEE Transactions on Software Engineering*,2013:39(6):757-773.
7. Hall T, Beecham S, Bowes D, Gray D, Counsell S. A systematic literature review on fault prediction performance in software engineering. *IEEE Transactions on Software Engineering*,2012:38(6):1276-1304.
8. da Silva NFF, de Souza EFG, de Oliveira CAS. Software defect prediction using machine learning techniques: A systematic review. *Journal of Systems and Software*, 2022, 185.
9. Witten IH, Frank E, Hall MA, Pal CJ. *Data Mining: Practical Machine Learning Tools and Techniques*, 4th ed. Cambridge, MA, USA: Morgan Kaufmann, 2016.
10. Pressman RS, Maxim BR. *Software Engineering: A Practitioner's Approach*, 9th ed. New York, NY, USA: McGraw-Hill, 2019.