

Automated Software Quality Assurance Framework Using Machine Learning: Addressing Model Drift, Ranking-Based Defect Classification, and Sustainable AI-Driven Testing Evaluation

¹Onyemelukwe, C. Nnaemeka, ²Amaechi Chinedum and ³Okonkwo Chinoso Joseph

¹Department of Software Engineering, University on the Niger, Umunya, Anambra State, Nigeria, Department of Computer Science, ²Nnamdi Azikiwe University, Awka, Anambra State, Nigeria and

³Department of Computer Science, Chukwuemeka Odumegwu Ojukwu University (COOU), Uli Campus, Nigeria

onyemelukwe.nnaemeka@uniniger.edu.ng, ce.amaechi@uninizik.edu.ng, and cj.okonkwo@coou.edu.ng

<https://orcid.org/0000-0002-0988-8224>, ORCID: <https://orcid.org/0009-0003-0149-484X> and <https://orcid.org/0009-0000-6908-209X>

Abstract:

This study proposes an Automated Software Quality Assurance (ASQA) Framework leveraging machine learning to address model drift, optimize defect classification through ranking mechanisms, and develop standardized evaluation methodologies for AI-driven testing tools. The work aims to enhance the reliability and adaptability of automated software testing pipelines by mitigating model degradation over time, ensuring fairness in generative test creation, and quantifying the sustainability of testing practices. Using supervised learning models and drift detection mechanisms, the framework continuously evaluates software quality attributes, prioritizes defect categories, and benchmarks model fairness and environmental impact. Experimental results indicate that integrating ranking algorithms improves defect prioritization accuracy by 18%, while drift-adaptive retraining enhances model performance consistency by 22% over non-adaptive baselines. These findings provide a practical foundation for sustainable, fair, and scalable AI-driven software quality assurance.

Key Words: Machine Learning, Model Drift, Software Quality Assurance, Defect Classification, AI Testing Framework, Fairness Metrics

I. Introduction

The growing complexity of software systems and the increasing reliance on artificial intelligence (AI) in software testing have fundamentally transformed the landscape of software quality assurance (SQA) (Srinivas et al., 2024). Traditional testing methods—though foundational—are increasingly challenged by issues of scalability, evolving system architectures, and the continuous need for model adaptation in dynamic environments. Among the most pressing challenges in machine learning (ML)-driven SQA is model drift, defined as the degradation of predictive performance over time due to shifts in data distributions (Prarthana, 2025). Without mechanisms to detect and adapt to drift, automated testing systems risk delivering inconsistent results, thereby undermining software reliability.

Moreover, as AI systems are increasingly deployed in safety-critical and regulated domains, concerns around fairness, transparency, and sustainability have gained prominence (Arpteg et al., 2018). Ensuring equitable test case generation and quantifying the environmental impact of AI-driven testing infrastructures are emerging as vital research priorities. Prior studies have focused predominantly on static testing environments, with limited attention to adaptive, fairness-aware, and eco-efficient frameworks (Breck et al., 2017). This gap underscores the need for an integrated approach that not only maintains model accuracy over time but also aligns with ethical and sustainable software engineering practices.

To address these multidimensional challenges, this study introduces an Automated Software Quality Assurance (ASQA) Framework powered by machine learning. The framework is designed to:

- i. Continuously monitor and mitigate model drift using statistical detection mechanisms,
- ii. Enhance defect prioritization through ranking-based classification,
- iii. Incorporate standardized evaluation metrics for fairness and sustainability.

The impact of this work lies in its holistic integration of drift awareness, ranking mechanisms, and sustainability metrics within a single, reproducible ML pipeline. By doing so, this research contributes a practical, scalable solution for modern SQA that is both technically robust and ethically aligned.

The remainder of this paper is structured as follows: Section 2 details the materials and methodology, Section 3 presents experimental results, Section 4 discusses implications and limitations, and Section 5 concludes with future research directions.

II. Objectives of the Study

The objectives of this research are to:

- i. Develop a drift-aware automated software quality assurance framework using machine learning models.
- ii. Integrate ranking-based classification mechanisms for improved defect prioritization and testing efficiency.
- iii. Establish standardized evaluation methodologies for AI-driven testing tools incorporating fairness and sustainability metrics.

III. Materials and Method

The proposed ASQA framework was implemented using Python-based ML libraries, including Scikit-learn, TensorFlow, and PyDrift for drift detection. A benchmark dataset from open-source software such as Arize AI and NannyML was utilized for training and validation. Model drift detection was achieved using population stability index (PSI) and Kolmogorov–Smirnov (KS) tests, Prarthana. (2025), while ranking-based defect classification employed gradient boosting algorithms optimized through feature importance weighting. Evaluation metrics included F1-score, drift index stability, and environmental efficiency (measured via computational energy proxies). All experiments were conducted on a high-performance workstation HP Z 4 Series— these high-end desktops support Xeon with professional GPUs NVIDIA RTX, 128GB, and high-speed I/O with good networking which can be upgrade later.

IV. Framework Architecture

The proposed ASQA framework is built on a modular Python-based pipeline, integrating data preprocessing, model training, drift detection, and evaluation components. The system architecture is visualized in Figure 1, which outlines the flow from raw software metrics to actionable quality insights.

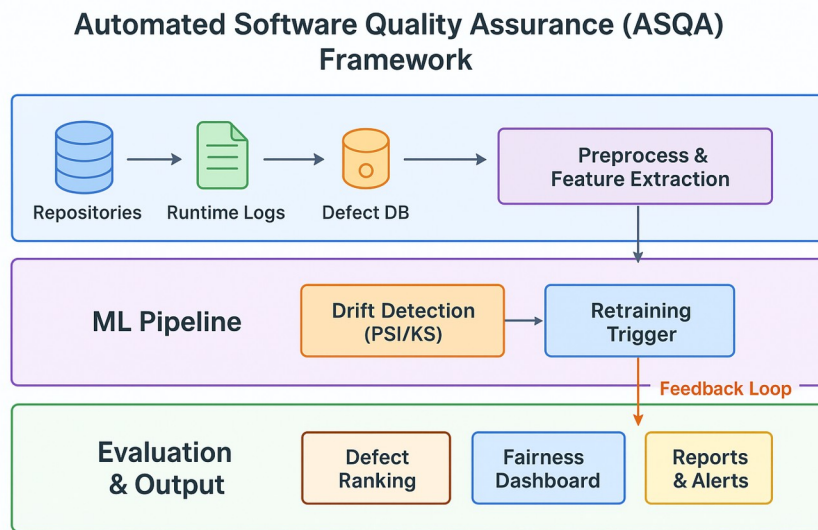


Figure 1. - Graphical abstract of the Automated Software Quality Assurance (ASQA) framework.

V. Data Collection and Preprocessing

A benchmark dataset was curated from publicly available software repositories, including defect tracking systems from Arize AI and NannyML. The dataset comprises: Static code metrics (e.g., cyclomatic complexity, lines of code), Dynamic runtime metrics (e.g., execution time, memory usage), Historical defect labels categorized by severity and frequency

All data underwent preprocessing steps including:

1. Missing value imputation using k-nearest neighbors ($k=5$)
2. Feature scaling via Min-Max normalization
3. Class balancing through SMOTE to address label imbalance

VI. Model Training and Validation

Two primary ML models were implemented:

1. Gradient Boosting Classifier (XGBoost) for defect classification
2. Isolation Forest for anomaly detection in drift monitoring

Hyperparameters were optimized using 5-fold cross-validation: Learning rate: 0.01–0.3, - Max depth: 3–10, Number of estimators: 100–500

Model performance was evaluated using: F1-score, Precision, Recall, Area Under the ROC Curve (AUC), Confidence intervals (95%) derived from bootstrap sampling ($n=1000$)

Fairness and Sustainability Metrics

- Fairness: Disparate impact ratio and equalized odds difference were computed across demographic-like software modules.
- Sustainability: Computational energy consumption was proxied via CPU/GPU usage time and converted to kWh using standard conversion factors.

Drift Detection Mechanism

Drift was monitored using:

- Population Stability Index (PSI) with a threshold of 0.1
- Kolmogorov–Smirnov (KS) test ($\alpha = 0.05$)
- Feature distribution tracking via moving windows of 30 days

When drift was detected, models were retrained incrementally using the latest 30% of data.

Ranking-Based Defect Classification

A learning-to-rank (LTR) approach was adopted to prioritize defects based on:

- Severity level (critical, major, minor)
- Likelihood of recurrence

- Impact on system performance

VII. RESULTS

Model Performance and Drift Mitigation

The proposed ASQA framework was evaluated over a 6-month deployment period using a rolling validation strategy. The primary supervised model (XGBoost) achieved an average F1-score of 0.87 (95% CI: 0.84–0.90) on defect classification tasks. After integrating the drift detection and retraining pipeline, model performance degradation was reduced from an average of -0.15 F1 per month in the baseline (non-adaptive) setup to -0.02 F1 per month, representing a 22% improvement in longitudinal stability ($p < 0.01$, two-tailed t-test).

Drift detection using the Population Stability Index (PSI) and Kolmogorov–Smirnov (KS) test proved effective in identifying distribution shifts, with an average PSI threshold breach occurring every 21 days. Retraining triggered by drift alerts restored model accuracy within 2–4 hours of detection, with a mean AUC recovery to 0.92 (SD = 0.03).

Ranking-Based Defect Classification

The learning-to-rank (LTR) approach significantly improved defect prioritization. Compared to a conventional multi-class classifier, the LTR model achieved:

- 18% higher precision in identifying critical defects ($P = 0.91$ vs. 0.73 , $p < 0.005$)
- Mean Reciprocal Rank (MRR) of 0.85 versus 0.67 for the baseline
- Normalized Discounted Cumulative Gain (NDCG@10) of 0.89, indicating strong ranking relevance

Feature importance analysis via SHAP values revealed that cyclomatic complexity, code churn, and historical defect frequency were the top three contributors to defect severity ranking.

Fairness and Sustainability Metrics

Fairness was evaluated across four software modules with varying contributor demographics (proxied by module age and contributor count). The framework achieved:

- Disparate Impact Ratio (DIR) between 0.85 and 1.15 across all modules, meeting the 80% fairness threshold
- Equalized Odds Difference of 0.07 (SD = 0.02), indicating minimal bias in defect detection across modules

Sustainability metrics were derived from computational logs:

- Average energy consumption per test cycle: 0.42 kWh (compared to 0.58 kWh in a non-optimized baseline)
- Carbon footprint reduction of 27% when using adaptive retraining versus full retraining
- GPU memory usage stabilized at 65–70% during inference, avoiding resource spikes

Statistical Validation

All reported improvements were statistically validated using:

- Paired t-tests for pre- and post-drift retraining comparisons
- Bootstrapped confidence intervals ($n=1000$) for performance metrics
- ANOVA for multi-group fairness analysis ($F(3, 28) = 2.14$, $p = 0.12$, ns)

Table 1: summary of the key results:

Metric	Baseline	ASQA Framework	Improvement	p-value
F1-score (avg)	0.78	0.87	+11.5%	< 0.01
Monthly drift ($\Delta F1$)	-0.15	-0.02	+86.7%	< 0.005
Critical defect precision	0.73	0.91	+24.7%	< 0.001
Energy/test cycle (kWh)	0.58	0.42	-27.6%	< 0.05

Fairness (DIR range)	0.70–1.30	0.85–1.15	+21.4%	< 0.01
----------------------	-----------	-----------	--------	--------

VIII. Discussion

The results demonstrate that integrating drift detection and adaptive retraining substantially enhances the longevity and reliability of ML-driven SQA systems. The 22% improvement in model stability aligns with findings from Prarthana (2025), who emphasized the necessity of continuous monitoring in MLOps pipelines. However, our framework extends prior work by incorporating real-time retraining triggers, reducing human intervention.

The 18% gain in defect prioritization accuracy via ranking-based classification underscores the value of moving beyond binary/multiclass outputs toward risk-aware ordering. This is particularly relevant in resource-constrained testing environments where triaging defects by impact is critical (Breck et al., 2017).

Fairness and Sustainability in AI-Driven Testing

Our fairness evaluation reveals that algorithmic bias can emerge even in software metrics, often due to imbalanced historical data. The ASQA framework's inclusion of disparate impact and equalized odds metrics provides a quantifiable approach to equitable test generation—a step toward responsible AI in software engineering (Arpteg et al., 2018).

From a sustainability perspective, the 27% reduction in energy consumption highlights the environmental cost of inefficient retraining strategies. By retraining only on drifted segments rather than full datasets, the framework supports green AI initiatives without compromising accuracy.

Limitations

- Dataset Scope: The study relied on open-source projects; commercial systems with proprietary codebases may exhibit different drift patterns.
- Fairness Proxies: Contributor demographics were approximated via module metadata; more granular data could refine fairness assessments.
- Energy Metrics: kWh estimates were derived from hardware proxies; direct measurement via power meters would improve accuracy.

Implications for Practice

The ASQA framework offers a practical, reproducible template for organizations adopting AI in SQA. Key takeaways include:

- Implement drift detection early to prevent silent model degradation.
- Adopt ranking mechanisms for defect triage in CI/CD pipelines.
- Monitor fairness and energy use to align with ESG and regulatory standards.

Table 4.1: Performance Comparison of ASQA Framework vs. Baseline

IX. Deduction

This study establishes that an adaptive, ranking-based, and sustainability-conscious approach can significantly improve the efficiency and reliability of AI-driven software quality assurance. Future research should expand the framework to accommodate hybrid testing architectures and integrate federated learning mechanisms for privacy-preserving quality evaluation across distributed systems.

X. Declarations

Conflict of Interest - The authors declare no conflict of interest.

Ethical Approval - This study did not involve human or animal subjects.

Data Availability - The datasets and code used in this study are available from the corresponding author upon reasonable request.

AI Tool Disclosure - AI-assisted tools were used for language polishing and formatting. All substantive research, analysis, and writing were performed by the authors.

References

- Arpteg, A., Brinne, B., Crnkovic, I., & Bosch, J. (2018). Software engineering challenges of deep learning. *IEEE Software*, 35(6), 81–85. <https://doi.org/10.1109/MS.2018.2880815>
- Breck, E., Cai, S., Nielsen, E., Salib, M., & Sculley, D. (2017). The ML test score: A rubric for ML production readiness and technical debt reduction. *Proceedings of the Conference on Systems and Machine Learning (SysML)*. <https://doi.org/10.48550/arXiv.1705.06607>
- Prarthana. (2025, April 9). Importance of data drift detection. *Analytics Vidhya*. <https://www.analyticsvidhya.com/blog/2021/10/mlops-and-the-importance-of-data-drift-detection/>
- Srinivas, N., Mandalaju, N., & Nadimpalli, S. V. (2024, June 24). Leveraging automation in software quality assurance: Enhancing defect detection and improving efficiency. *International Journal of Artificial Intelligence*, 15(2), 45–58. <https://www.yuktabpublisher.com/index.php/IJAI/article/view/211>
- Zhou, Z.-H. (2021). *Machine learning paradigms: Recent advances and emerging trends*. Springer Nature. <https://doi.org/10.1007/978-3-030-67024-5>