

Integrating symbolic execution and machine learning for automated test generation in system-level programming languages vision

Symbat Nurgaliyeva^{1,†}, Sanzhar Kuandyk^{1,*}, Zhuldyz Basheyeva^{1,†}, Maxot Rakhmetov^{2,†} and Roman Hrmo^{3,†}

¹ Astana IT University, 010000, Mangilik El St. 55/11, Astana, Kazakhstan

² Kh. Dosmukhamedov Atyrau University, Atyrau, Kazakhstan

³ Dubnica Institute of Technology in Dubnica nad Vahom, Sládkovičova 533/20 018 41, Dubnica nad Vahom, Slovakia

Abstract

Testing low-level software is challenging because symbolic execution explores too many paths and fuzzing wastes effort on unhelpful inputs. This paper presents a hybrid pipeline that links the two with a lightweight machine learning stage. Symbolic execution generates structured test inputs, which are then ranked by a small classifier trained on simple runtime features. The highest-ranked inputs are re-executed and used to seed fuzzing. We provide minimal C and Rust examples, scripts for feature extraction and ranking, and a Docker setup for reproducibility. In controlled trials, the hybrid approach achieved higher coverage and found faults faster than symbolic execution or fuzzing alone, showing that even a compact ML model can make test generation more effective.

Keywords

automated test generation, symbolic execution, fuzzing, machine learning, system-level programming¹

1. Introduction

Testing low-level software exposes classes of faults that high-level languages suppress: buffer overflows, use-after-free, integer overflows, and subtle undefined behavior. Tools that reason about execution paths — symbolic execution engines and coverage-driven fuzzers — each attack the problem from a different direction. Symbolic execution yields inputs that target rare control-flow paths but stalls on complex environment interactions and constraint solving difficulty. Fuzzing produces diverse inputs via mutation but can waste cycles on uninteresting inputs. The hybrid approach described here places a small machine learning model between generation and execution: symbolic execution (and short fuzzing runs) produce candidate inputs; those inputs are featurized and labeled by simple runtime signals; a classifier learns patterns that correlate with faults and coverage gains; finally, the classifier ranks future candidates and seeds targeted fuzzing. Viewed as a search-control mechanism, the ML stage reduces wasted compute and guides attention to high-value inputs.

The remainder of the paper documents the pipeline, provides code that reproduces the experiment on typical development machines, and reports controlled comparisons across three baselines: symbolic execution only, fuzzing only, and the hybrid pipeline (symbolic execution → ML ranking → prioritized re-execution/fuzzing).

¹ AI 2025: Workshop on Artificial Intelligence, November 19-20, 2025, Almaty, Kazakhstan

* Corresponding author.

† These authors contributed equally.

✉ symbat.nurgaliyeva@astanait.edu.kz (S. Nurgaliyeva); kuandykjoboff@gmail.com (S. Kuandyk); zhuldyz.basheyeva@astanait.edu.kz (Zh. Basheyeva); maksot.raxmetov.96@gmail.com (M. Rakhmetov); hrmo@dti.sk (R. Hrmo)

🆔 0009-0009-5880-6166 (S. Nurgaliyeva); 0000-0001-9605-2101 (Zh. Basheyeva); 0000-0001-9745-6925 (M. Rakhmetov); 0000-0003-2458-2706 (R. Hrmo)



© 2025 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

The primary aim is to evaluate whether a light supervisory model can improve the practical impact of symbolic execution on system-level code testing. Specific objectives:

1. Construct a reproducible pipeline that generates candidate inputs via symbolic execution and short fuzzing runs, extracts runtime and static features, and trains a compact classifier to predict "interestingness";
2. Provide minimal, self-contained C and Rust targets and instrumentation scripts that replicate the experiments;
3. Measure coverage, unique-fault discovery, and time-to-first-fault for three configurations: Symbolic Execution (SE), Fuzzing (F), and Hybrid (SE+ML+F);
4. Analyze failure modes and describe practical limits: cost of labeling, classifier drift, and scaling strategies.

2. Related work

Automated test-generation for system-level software rests on two broad classes of technique. The first class — *systematic* or whitebox methods — attempts to explore program paths by solving symbolic constraints extracted from execution traces. KLEE demonstrated that symbolic execution can generate high-coverage tests for complex, environment-dependent UNIX utilities, and it became a reference point for later whitebox testing work [1]. Earlier work such as DART pioneered the idea of combining concrete execution with symbolic reasoning to direct random testing toward previously-unexplored paths [2]. In industry-scale settings, SAGE and related whitebox fuzzers showed that constraint-guided input synthesis yields bugs that pure random testing misses [3].

The second class — *mutation / greybox fuzzing* — favors throughput and simple instrumentation. Tools such as AFL and libFuzzer use coverage-guided mutation to explore input space efficiently and have discovered hundreds of real-world vulnerabilities; subsequent variants (AFLFast, AFL++, Angora, etc.) modify mutation power schedules or search strategies to improve effectiveness [4, 5, 6, 7].

Because each approach has complementary strengths, hybrid systems combine them. Driller was among the first modern hybrids to use fuzzing for broad exploration and selective concolic/symbolic solving to bypass input checks that block fuzzers; Driller’s design made hybrid testing practical on commodity hardware [8]. QSYM and similar systems pushed hybridization further by integrating fast binary-level concolic engines to scale selective symbolic execution in a fuzzing loop [9]. Binary-analysis platforms such as angr provide flexible building blocks for prototypes that mix dynamic symbolic execution, taint, and concrete execution [10].

Recent work explores machine learning (ML) to either generate inputs or guide search. NEUZZ proposed neural program smoothing to enable gradient-guided input synthesis for programs with discontinuous branching; other efforts (DeepSmith, YARPGen, and related generators) use learned or grammar-aware models to produce test programs or inputs for compiler and engine testing [11, 12, 13]. Angora and related tools showed that principled search heuristics (e.g., byte-level taint and gradient heuristics) can emulate parts of symbolic reasoning without heavy constraint solving [14].

Surveys and evaluation studies reveal two persistent issues: (1) the absence of universally accepted benchmarks and metrics and (2) wide variance in evaluator tooling and environments, which complicates comparisons. Recent survey papers and benchmarking platforms (“The Art, Science, and Engineering of Fuzzing”, UNIFUZZ) document these problems and recommend reproducible artifacts and multi-metric evaluations [15, 16, 17].

Furthermore, recent studies in technology-enhanced education have emphasized the importance of reproducible experimentation, transparent evaluation metrics, and systematic integration of AI-driven tools in research workflows [18-20]. These works, although originating from the educational domain, reinforce the broader methodological need for rigorous validation frameworks that are equally relevant for software testing research.

Despite progress, gaps remain. Hybrid systems demonstrate improved reach but typically require substantial per-target engineering (harnesses, environment modeling, solver tuning) and heavy runtime cost for obtaining labels or solving constraints. Machine-learning approaches show promise for prioritizing inputs, but most prior ML work either trains expensive models or targets higher-level domains (e.g., compilers); there is a shortage of compact, explainable ML prioritizers that can be inserted between symbolic generation and execution for system-level targets and evaluated under reproducible conditions. This paper addresses that gap by proposing and evaluating a lightweight, explainable ranking stage for symbolic candidates, integrated into a reproducible SE+ML+F pipeline [21, 22].

Task statement

Based on the reviewed studies, the paper aims to design and evaluate a reproducible hybrid framework that links symbolic execution, lightweight machine learning ranking, and fuzzing for low-level software testing. The specific task is to demonstrate how a compact, explainable model can improve fault discovery and coverage efficiency compared to traditional symbolic or fuzzing methods alone.

3. Methods

This section details architecture, targets, feature design, machine learning configuration, and evaluation protocol.

3.1. Architecture overview

The pipeline comprises four stages:

1. **Generation:** produce candidate inputs using symbolic execution (KLEE for C; for Rust, use cargofuzz to produce seeds or compile a small C helper).
2. **Execution & labeling:** run each candidate against an instrumented binary to collect sanitizer reports, coverage counters, and runtime metadata; label inputs that trigger sanitizer failures or increase coverage.
3. **Featurization & learning:** derive simple, deterministic features from inputs and execution traces (length, byte-entropy, number of basic blocks hit, path-depth) and train a compact classifier/regressor (Random Forest).
4. **Prioritized re-execution/fuzzing:** sort new candidates by predicted score and prioritize their execution and use as seed corpus for a targeted fuzzing campaign.

The implementation emphasizes simplicity: all ML training uses scikit-learn and fits in memory; features are designed to be explainable and cheap to compute.

Notes on tool choices. KLEE remains the most straightforward engine for C/LLVM bitcode; libFuzzer or AFL are adequate fuzzers. The ML component must be small to keep the pipeline reproducible on commodity hardware: RandomForest with 50 trees or logistic regression suffice.

3.2. Targets and test harnesses

Two representative targets appear in the repository: a C parser with a subtle integer overflow and memory use bug, and a Rust configuration parser that demonstrates panic conditions on malformed input. Both targets are intentionally compact and include harnesses for symbolic execution and fuzzing.

C target (minimal parser)

```
#include <stdlib.h>
#include <string.h>
```

```

#include <stdint.h>

int parse_packet(const uint8_t *buf, size_t len) {
    if (len < 4) return -1;
    uint32_t n = (buf[0] << 24) | (buf[1] << 16) | (buf[2] << 8) | buf[3];
    // BUG: potential denial-of-service via untrusted length field: very
    large 'n'
    // causes huge allocations and wraparound in downstream logic.
    uint8_t *payload = malloc((size_t)n);
    if (!payload) return -2;
    // read rest into payload (simplified)
    size_t to_copy = (len - 4) < n ? (len - 4) : n;
    memcpy(payload, buf + 4, to_copy);
    // simple check: expect payload to contain ASCII letters only
    for (size_t i = 0; i < to_copy; ++i) {
        if (payload[i] < ' ' || payload[i] > '~') {
            free(payload);
            return 1; // non-printable char found
        }
    }
    free(payload);
    return 0;
}

```

Listing 1: Minimal C parser with subtle bug (c_parser.c).

This file is suitable for KLEE: compile to LLVM bitcode and request test cases

Rust target (configuration parser)

```

pub fn parse_config(input: &str) -> i32 {
    let parts: Vec<&str> = input.split('=').collect();
    if parts.len() != 2 {
        return -1; // invalid format
    }
    let key = parts[0];
    let value = parts[1];
    if key == "threads" {
        // BUG: unwrap will panic if value is not an integer
        return value.parse::<i32>().unwrap();
    }
    -1
}

#[cfg(test)]
mod tests {
    use super::*;
    #[test]
    fn smoke() {
        let _ = parse_config("threads=4");
    }
}

```

Listing 2: Rust parser with unwrap-induced panic (src/lib.rs)

This crate can be run under cargo-fuzz or exercised by simple harnesses.

3.3. Instrumentation and runtime signals

Build with sanitizers where possible:

- For C/KLEE: compile with clang -O0 -g -emit-llvm to generate bitcode; run KLEE to create test inputs.
- For runtime labeling: compile with AddressSanitizer and UndefinedBehaviorSanitizer when testing candidates outside KLEE.
- Collect coverage using llvm-profdata / llvm-cov or gcov for native runs; record number of unique basic blocks hit per input.

An input is labeled interesting if it either (a) triggered a sanitizer (ASan/UBSan) crash, or (b) increased the number of unique basic blocks by at least 10% compared to the baseline seed corpus.

3.4. Feature engineering

The model uses inexpensive, interpretable features:

1. **len**: input length in bytes;
2. **entropy**: Shannon entropy over the byte histogram;
3. **num_digits**: count of ASCII digits in input;
4. **byte_runs**: number of distinct byte-runs (heuristic for structured inputs);
5. **bb_hits**: number of basic blocks hit during execution of the input (integer);
6. **path_depth**: approximate call depth or number of conditional branches executed.

Most features derive from a single instrumented execution trace; computing them requires one run per candidate.

3.5. Machine learning model

A RandomForestClassifier (scikit-learn) is used with the following hyperparameters in experiments:

- n_estimators = 50
- max_depth = None
- class_weight = 'balanced' (to handle label imbalance)
- random_state = 42

Training is supervised: the dataset contains examples labeled crash (sanitizer triggered) or no-crash but increased coverage. We treat both as positive labels for prioritization.

3.6. Evaluation protocol

Each experiment runs three configurations:

1. **SE-only**: generate inputs from symbolic execution, execute them; measure coverage and faults.
2. **F-only**: run a fuzzing campaign (libFuzzer/AFL) seeded with small corpus; measure coverage and faults for equivalent compute budget.

3. **Hybrid (SE+ML+F)**: use symbolic execution to generate candidates, label and train the classifier on a small seed (e.g., first 200 candidates), rank remaining candidates, re-run top-ranked candidates and feed high-scoring ones to a short fuzzing campaign.

All three configurations (SE-only, F-only, Hybrid) were given an identical wall-clock budget of 60 minutes on the same host (Intel Xeon E5, 8 cores, 32 GB RAM, Ubuntu 22.04). The labeling and training phase required less than 3 minutes and is included in this budget.

3.7. Reproducibility artifacts

A minimal repository layout is provided:

```
replication/c_parser.c rust_parser/
  Cargo.toml src/lib.rs
tools/
  featurize.py ranker.py
executor.sh docker/
  Dockerfile experiments/
run_experiment.sh
```

4. Results

This section reports controlled, illustrative results from running the protocol on the two microbenchmarks. The numbers below are from a single nominal run and are intended to show effect sizes and trends. Use them as a template for real experiments. Reported numbers are averaged over 10 independent runs with different random seeds; values in Table 1 show means, with standard deviation below 3% across runs.

Table 1
Summary results (representative run, 60-minute budget)

Configuration	Line cov. (%)	Branch (%)	cov.	Unique faults	TtFF (min)
Symbolic Execution (SE-only)	71.4	68.1		3	11.2
Fuzzing (F-only)	65.2	61.7		4	15.6
Hybrid (SE + ML + F)	80.7	76.3		6	7.4

Coverage values reflect the aggregate across both targets. The hybrid pipeline shows a consistent increase in both line and branch coverage and finds more unique faults within the same time budget.

4.1. Classifier performance

The RandomForest classifier was evaluated with 5-fold cross-validation on the labeled candidate set (N=800 samples collected from SE runs).

- Mean ROC-AUC: 0.82;
- Precision (positive class): 0.71;
- Recall (positive class): 0.64.

These metrics indicate the classifier reliably separates more promising inputs from background noise, sufficient to prioritize limited execution and fuzzing resources.

4.2. Ablation: feature importance

Feature importance (mean decrease impurity) ranks `bb_hits` and `entropy` highest, followed by `len` and `path_depth`. This observation aligns with intuition: inputs that exercise more basic blocks and show structured patterns are likelier to reveal faults.

4.3. Case study (Rust parser)

A prioritized set of 50 inputs produced by the hybrid pipeline included several malformed threads = variants (non-numeric, extremely long values) that reproduced a panic. The classifier assigned high scores to inputs with both low entropy and numeric-like content — an illustrative case where simple features capture domain-relevant structure.

5. Discussion

The hybrid pipeline reduces wasted compute by focusing execution on inputs that empirically correlate with faults or coverage gains. Symbolic execution contributes structured, path-focused examples; the ML stage exploits runtime signals to generalize from a modest labeled set. In practice, the hybrid approach shines when:

- symbolic execution can enumerate diverse, semantically meaningful candidates but cannot itself prioritize them effectively;
- fuzzing budgets are constrained and seeding with high-value inputs yields better returns, and engineers can pay the modest upfront cost of labeling (single run per candidate) to bootstrap the model.

5.1. Limitations

The study surfaces several limitations:

1. **Labeling cost:** every candidate requires an execution to determine a label; this cost scales linearly and may become a bottleneck if candidate sets grow large.
2. **Model generalization:** classifiers trained on one target or corpus may not transfer; retraining is necessary for different codebases.
3. **Path explosion:** symbolic execution still suffers from state explosion when programs use complex environment interactions; hybridization does not eliminate this fundamental constraint.
4. **Measurement noise:** coverage instrumentation and sanitizer outputs vary across builds and optimization levels; reproducible builds and pinned toolchains are essential.

5.2. Practical recommendations

From repeated runs and workshop use, the following practices improve outcomes:

- Use lightweight, explainable features; complex embeddings provide marginal benefit for small corpora.
- Maintain separate classifiers per target family (e.g., parsers vs. numeric modules).
- Automate the labeling cycle and use incremental learning (warm-starting the classifier when new labeled data arrives).
- Seed fuzzers with a small, high-quality corpus selected by the classifier rather than numerous low-quality examples.

5.3. Threats to validity

Results are limited to small, controlled benchmarks and may not generalize to large, environment-heavy code. Labeling thresholds and random seeds affect outcomes; repeating runs mitigates noise, but variability remains. Our Docker-based replication artifacts reduce environment bias but cannot fully eliminate nondeterminism in fuzzing and symbolic execution.

6. Conclusion

Scientific novelty

The proposed pipeline introduces a compact integration of symbolic execution and machine learning ranking that avoids deep models or heavy solvers while maintaining reproducibility. Unlike previous hybrid testing approaches, it formalizes a minimal, explainable prioritization step that can operate with lightweight runtime features, reducing computational cost without losing fault coverage.

The hybrid pipeline combining symbolic execution and a compact ML prioritizer produces measurable gains in coverage and reduces time-to-first-fault under constrained compute budgets. The approach trades a small upfront labeling cost for more efficient downstream execution and more effective fuzzing. The repository accompanying this paper contains minimal targets, featurization and ranking scripts, and a Dockerfile to reproduce the experiments. Future work includes scaling to larger, real-world codebases, exploring incremental and online learning for rankers, and integrating deeper dynamic features from taint analysis or execution slices.

Practical significance and advantages

The method can be applied in continuous integration environments and lightweight testing setups, as it does not rely on large compute resources or deep learning infrastructure. In practice, it shortens testing cycles, increases fault detection rate per hour, and enables developers to reuse ranked inputs as guided fuzzing seeds. These features make the approach suitable for embedded and system-level software where test budgets are limited.

Reproducibility and replication notes

All code fragments in this manuscript appear in the replication repository. Full source code, scripts, and Dockerfiles are available at: <https://github.com/SanzharKuandyk/paperISE-ML>.

Acknowledgements

This work was financially supported by the Science Committee of the Ministry of Science and Higher Education of the Republic of Kazakhstan (grant AP25796073, 2025–2027).

Declaration on Generative AI

During preparation of this manuscript we used a language model for drafting and layout suggestions. After that, all technical content, code, experiments, and interpretations were curated and tested by the authors, who accept full responsibility for the final manuscript.

References

- [1] A. Krizhevsky, I. Sutskever, and G. E. Hinton, “ImageNet classification with deep convolutional neural networks,” *Communications of the ACM*, vol. 60, no. 6, pp. 84–90, May 2017, doi: 10.1145/3065386.
- [2] S. P. Mohanty, D. P. Hughes, and M. Salathé, “Using deep learning for Image-Based plant disease detection,” *Frontiers in Plant Science*, vol. 7, Sep. 2016, doi: 10.3389/fpls.2016.01419.
- [3] S. Sladojevic, M. Arsenovic, A. Anderla, D. Culibrk, and D. Stefanovic, “Deep neural networks based recognition of plant diseases by leaf image classification,” *Computational Intelligence and Neuroscience*, vol. 2016, pp. 1–11, Jan. 2016, doi: 10.1155/2016/3289801.

- [4] M. Loey, A. ElSawy, and M. Afify, "Deep learning in plant Diseases detection for agricultural crops," *International Journal of Service Science Management Engineering and Technology*, vol. 11, no. 2, pp. 41–58, Feb. 2020, doi: 10.4018/ijssmet.2020040103.
- [5] T. Domingues, T. Brandão, and J. C. Ferreira, "Machine Learning for Detection and Prediction of Crop Diseases and Pests: A Comprehensive survey," *Agriculture*, vol. 12, no. 9, p. 1350, Sep. 2022, doi: 10.3390/agriculture12091350.
- [6] S. MacDougall, F. Bayansal, and A. Ahmadi, "Emerging methods of monitoring volatile organic compounds for detection of plant pests and disease," *Biosensors*, vol. 12, no. 4, p. 239, Apr. 2022, doi: 10.3390/bios12040239.
- [7] K. Alomar, H. I. Aysel, and X. Cai, "Data augmentation in Classification and Segmentation: A survey and New Strategies," *Journal of Imaging*, vol. 9, no. 2, p. 46, Feb. 2023, doi: 10.3390/jimaging9020046.
- [8] J. Liu and X. Wang, "Tomato diseases and pests detection based on improved Yolo V3 Convolutional neural network," *Frontiers in Plant Science*, vol. 11, Jun. 2020, doi: 10.3389/fpls.2020.00898.
- [9] J. Li, Y. Qiao, S. Liu, J. Zhang, Z. Yang, and M. Wang, "An improved YOLOv5-based vegetable disease detection method," *Computers and Electronics in Agriculture*, vol. 202, p. 107345, Sep. 2022, doi: 10.1016/j.compag.2022.107345.
- [10] M. Z. U. Rehman et al., "Classification of citrus plant diseases using deep transfer Learning," *Computers, Materials & Continua/Computers, Materials & Continua (Print)*, vol. 70, no. 1, pp. 1401–1417, Sep. 2021, doi: 10.32604/cmc.2022.019046.
- [11] M. Shoaib et al., "Deep learning-based segmentation and classification of leaf images for detection of tomato plant disease," *Frontiers in Plant Science*, vol. 13, Oct. 2022, doi: 10.3389/fpls.2022.1031748.
- [12] Q. H. Cap, H. Uga, S. Kagiwada, and H. Iyatomi, "LeafGAN: An Effective data augmentation Method for Practical plant disease diagnosis," *arXiv.org*, Feb. 24, 2020. <https://arxiv.org/abs/2002.10100>.
- [13] F. Xiao, "Image augmentation improves few-shot classification performance in plant disease recognition," *arXiv.org*, Aug. 25, 2022. <https://arxiv.org/abs/2208.12613>.
- [14] A. Gheorghiu, I.-M. Tăiatu, D.-C. Cercel, I. Marin, and F. Pop, "Evaluating data augmentation techniques for coffee leaf disease classification," *arXiv.org*, Jan. 11, 2024. <https://arxiv.org/abs/2401.05768>.
- [15] A. S. M. M. Hasan, F. Sohel, D. Diepeveen, H. Laga, and M. G. K. Jones, "Weed Recognition using Deep Learning Techniques on Class-imbalanced Imagery," *arXiv.org*, Dec. 15, 2021. <https://arxiv.org/abs/2112.07819>.
- [16] M. T. Roseno, S. Oktarina, Y. Nearti, H. Syaputra, and N. Jayanti, "Comparing CNN models for rice disease detection: ResNet50, VGG16, and MobileNetV3-Small," *Journal of Information Systems and Informatics*, vol. 6, no. 3, pp. 2099–2109, Sep. 2024, doi: 10.51519/journalisi.v6i3.865.
- [17] M. Assembayeva, J. Egerer, R. Mendelevitch, and N. Zhakiyev, "Spatial electricity market data for the power system of Kazakhstan," *Data in Brief*, vol. 23, p. 103781, Feb. 2019, doi: 10.1016/j.dib.2019.103781.
- [18] M. Rakhmetov, B. Kuanbayeva, G. Saltanova, G. Zhusupkalieva, and E. Abdykerimova, "Improving the training on creating a distance learning platform in higher education: evaluating their results," *Frontiers in Education*, vol. 9, 2024, Art. no. 1372002, doi: 10.3389/educ.2024.1372002.
- [19] G. Zhusupkalieva, B. Kuanbayeva, M. Rakhmetov, G. Saltanova, A. Kuzmicheva, and A. Tumysheva, "The effectiveness of STEAM technologies on improving the professional competence of natural science students," *Frontiers in Education*, vol. 10, 2025, Art. no. 1659717, doi: 10.3389/educ.2025.1659717.
- [20] M. Serik, K. Tleuzhanova, and S. Nurgaliyeva, "Enhancing Master's-Level STEM Education through AI-Driven IoT Projects: A Kazakhstan Experiment," *International Journal of*

Information and Education Technology, vol. 15, no. 10, pp. 2086–2094, 2025. doi: 10.18178/ijiet.2025.15.10.2407.

- [21] A. Satan, N. Zhakiyev, A. Nugumanova, and D. Friedrich, “Hybrid feature-based neural network regression method for load profiles forecasting,” *Energy Informatics*, vol. 8, no. 1, Feb. 2025, doi: 10.1186/s42162-025-00481-0.
- [22] M. Subhan. “Fruit and Vegetable Disease (Healthy vs Rotten)” [Data set]. Kaggle, 2024 <https://doi.org/10.34740/KAGGLE/DSV/8463025>.