

The Repeat Offenders: Characterizing and Predicting Extremely Bug-Prone Source Methods

ETHAN FRIESEN, SQM Research Lab, Computer Science, University of Manitoba, Canada

SASHA MORTON-SALMON, SQM Research Lab, Computer Science, University of Manitoba, Canada

MD NAHIDUL ISLAM OPU, SQM Research Lab, Computer Science, University of Manitoba, Canada

SHAHIDUL ISLAM, SQM Research Lab, Computer Science, University of Manitoba, Canada

SHAIFUL CHOWDHURY, SQM Research Lab, Computer Science, University of Manitoba, Canada

Identifying the small subset of source code that repeatedly attracts bugs is critical for reducing long-term maintenance effort. We define *ExtremelyBuggy* methods as those involved in more than one bug fix and present the first large-scale study of their prevalence, characteristics, and predictability. Using a dataset of over 1.25 million methods from 98 open-source Java projects, we find that *ExtremelyBuggy* methods constitute only a tiny fraction of all methods, yet frequently account for a disproportionately large share of bugs. At their inception, these methods are significantly larger, more complex, less readable, and less maintainable than both singly-buggy and non-buggy methods. However, despite these measurable differences, a comprehensive evaluation of five machine learning models shows that early prediction of *ExtremelyBuggy* methods remains highly unreliable due to data imbalance, project heterogeneity, and the fact that many bugs emerge through subsequent evolution rather than initial implementation. To complement these quantitative findings, we conduct a thematic analysis of 265 *ExtremelyBuggy* methods, revealing recurring visual issues (e.g., confusing control flow, poor readability), contextual roles (e.g., core logic, data transformation, external resource handling), and common defect patterns (e.g., faulty conditionals, fragile error handling, misuse of variables). These results highlight the need for richer, evolution-aware representations of code and provide actionable insights for practitioners seeking to prioritize high-risk methods early in the development lifecycle.

CCS Concepts: • **Software and its engineering** → *Maintaining software*.

Additional Key Words and Phrases: *ExtremelyBuggy*, software bug, bug-proneness, bug prediction, software maintenance, code metrics.

1 Introduction

Software maintenance is one of the most expensive phases of the development lifecycle [5]. A major contributor to this cost is the effort required to identify and correct software bugs, which alone can account for 50-70% of the total development costs [13, 73]. To no surprise, building models for accurate bug prediction is referred to as the *prince of empirical software engineering research* [37] with researchers developing bug prediction models in an attempt to take early actions [13, 20, 50, 56, 66, 73, 74]. The assumption being that if bug-prone code components can be identified early in the development lifecycle or prevented altogether, future maintenance costs will be significantly reduced [7].

Prior bug prediction models predominantly focused on the class or file levels [2, 21, 75]. Unfortunately, these models remain underutilized in practice [72], at least partly because practitioners find it challenging to locate bugs at such coarse levels of granularity [13, 24, 37, 43, 50, 64]. Practical use would require developers to sift through entire classes or files without an indication of what characteristics the model considered bug-prone. Consequently, a significant amount of recent research has focused on bug prediction at the method level granularity [13, 19, 20, 41, 50, 65]. However, these

Authors' Contact Information: Ethan Friesen, SQM Research Lab, Computer Science, and University of Manitoba, Winnipeg, Canada, fries432@myumanitoba.ca; Sasha Morton-Salmon, SQM Research Lab, Computer Science, and University of Manitoba, Winnipeg, Canada, mortonss@myumanitoba.ca; Md Nahidul Islam Opu, SQM Research Lab, Computer Science, and University of Manitoba, Winnipeg, Canada, opummi@myumanitoba.ca; Shahidul Islam, SQM Research Lab, Computer Science, and University of Manitoba, Winnipeg, Canada, islams32@myumanitoba.ca; Shaiful Chowdhury, SQM Research Lab, Computer Science, and University of Manitoba, Winnipeg, Canada, shaiful.chowdhury@umanitoba.ca.

models treat all buggy methods equally, without distinguishing between methods that were fixed once and those that required multiple bug fixes.

In contrast to previous research, we focus on *ExtremelyBuggy* methods—methods that have required bug fixes more than once—because identifying them early could prevent a disproportionate number of future failures and reduce the need for repeated maintenance on the same code fragments. Such methods may indicate architectural weak points or persistently complex areas that tend to accumulate defects over time. By identifying and targeting these methods, we aim to help practitioners prioritize their efforts where it matters most, enabling them to address a large number of bugs by focusing on a relatively small subset of problematic methods.

To support this goal, we analyze the characteristics of *ExtremelyBuggy* methods in an effort to predict them at their inception—that is, as soon as they are pushed to a software project. To do so, we collected the change history of 1.25 million methods from 98 widely used open-source Java projects. To the best of our knowledge, this represents the largest and most comprehensive dataset to date focused on method-level bug prediction. Using this dataset, we address the following four research questions:

RQ1: What proportion of methods are *ExtremelyBuggy*?

Contribution 1: We found that 0.04-6.63% of methods are *ExtremelyBuggy*. However, this small proportion of methods can often account for a large number of bug fixes in a project. This implies that identifying this small fraction of methods early will significantly reduce future maintenance burden.

RQ2: Do *ExtremelyBuggy* methods exhibit significantly different code quality than others?

Contribution 2: We found that, at their inception, *ExtremelyBuggy* methods are significantly larger, less readable, and have lower maintainability scores compared to both *Buggy* and *NotBuggy* methods. This is encouraging, as it suggests that these methods exhibit distinguishable code quality that could be leveraged by machine learning models for early prediction. That led us to our next research question.

RQ3: Can code quality-based machine learning models predict these *ExtremelyBuggy* methods at their inception?

Contribution 3: Unfortunately, our results indicate that machine learning algorithms perform poorly in distinguishing *ExtremelyBuggy* methods from other methods. This is partly due to a significant class imbalance, as *ExtremelyBuggy* methods represent only a small fraction of the overall dataset. Common techniques such as oversampling and under-sampling did not lead to meaningful improvements in prediction performance.

RQ4: What are the observable characteristics of *ExtremelyBuggy* methods?

Contribution 4: The futility in predicting with machine learning models led us to manually investigate the characteristics of the *ExtremelyBuggy* methods so that we can provide actionable guidance for the practitioners. As such, we conducted a thematic analysis on a curated dataset of 287 *ExtremelyBuggy* methods. In general, we focused on three different aspects while applying thematic analysis: visual issues, context of a method, and bug-fix type. Our analysis reveals some common themes that exist within these 287 methods. For example, we found that methods that deal with the core logic and algorithms tend to become *ExtremelyBuggy*.

To help replication and extension in future method-level bug prediction research, we share this dataset publicly.¹

2 Related Work

In this section, we discuss previous studies involving bug prediction and code metrics to show how they have motivated this paper.

¹<https://github.com/SQMLab/ExtremelyBuggyPublicData/tree/main/Dataset>

2.1 Software Maintenance and Bug Prediction

Software maintenance is expensive [5] and the bug-proneness of code components has been identified as one of the largest indicators of future maintenance burdens [9]. As such, there has been an enormous effort by the research community over the last forty years [38] to predict bug-prone code components so that early actions can be taken [9, 16]. Despite the great amount of research dedicated to predicting bugs, bug prediction models have yet to see widespread industry use [37, 72]. This is due, at least in part, to the fact that prior studies have predominantly focused on bug prediction at the file or class level [2, 21, 75], a granularity that practitioners often consider too coarse for practical use [13, 24, 43, 50].

These concerns have led to a recent shift in bug prediction research towards the method level [19, 20, 25] granularity. However, multiple recent studies [13, 50] have shown that the high prediction accuracy² reported by these studies is often inflated due to unrealistic evaluation scenarios, and, when tested on realistic scenarios, the accuracy of these models drops significantly. Additionally, prior studies have not differentiated between methods associated with a single bug fix and those that are repeatedly involved in bug-fixing. This distinction is critical, as methods that frequently require maintenance—often referred to as maintenance hotspots—should be prioritized and allocated additional resources.

Motivated by this important gap in prior research, this paper focuses exclusively on source code methods that have been associated with bugs more than once. We investigate whether these methods exhibit distinct characteristics and whether they can be identified early in their lifecycle.

2.2 Code Metrics

For bug prediction models to work, they need some quantifiable metrics that characterize the source code to use as inputs to the models. These metrics often come in the form of product metrics [11, 19, 41, 50] and process metrics [22, 41, 55]. Process metrics are typically derived from the change history of a code component and have been shown to be effective predictors of bug-prone components [20, 55]. However, such metrics are not available immediately after a piece of code is written and therefore cannot support the early identification of maintenance-prone components. Waiting for sufficient change history to accumulate may allow these methods to cause significant issues, leading to high maintenance costs and potentially eroding stakeholder confidence.

Product metrics, on the other hand, are often readily available and easy to compute at the inception of a code component. As a result, building maintenance prediction models using only product metrics has been a long-standing goal in software engineering research [9, 46, 50, 61]. Despite their extensive usage in software engineering research, some studies have found that product metrics, apart from *size*, are not helpful for building models that predict future maintenance burdens [17, 21, 67]. However, these studies with negative results were conducted at the class/file level, whereas studies that focused on method-level granularity have found product metrics as useful predictors of maintenance [11–13, 36].

Motivated by the need for early prediction and the practicality of method-level source code metrics, we focus exclusively on these metrics to characterize and predict ExtremelyBuggy methods.

3 Methodology

In this section, we outline the overall methodology, including project selection, change history tracking, and the identification of bug-proneness indicators. We also detail the data preprocessing steps and the selection of code

²Unless otherwise stated, accuracy in this paper refers to precision, recall, and F1-score.

metrics. To improve readability, methodologies specific to individual research questions are presented alongside each corresponding research question.

3.1 Dataset

A substantial portion of our dataset originates from the work of Chowdhury et al. [13], which comprises 774,051 Java methods drawn from 49 open-source projects. Each method in this dataset is accompanied by its complete change history and is enriched with metadata such as commit messages, timestamps, authorship information, and change types. The dataset additionally includes code diffs for all historical modifications, binary labels indicating whether each change corresponds to a bug fix, and the number of methods modified within the associated commit. These 49 projects have been widely utilized, either fully or partially, in prior method-level software engineering studies [9, 12, 13, 21, 24, 49, 58, 69]. Building upon this foundation, we extend the dataset by collecting additional open-source projects and extracting their method-level change histories.

3.1.1 Project Selection. We employed the GitHub REST API³ and the PyGitHub library⁴ to collect open-source projects. To ensure the quality and reliability of our dataset, we followed recommended best practices for mining GitHub data [35]. We began by retrieving the 1000 most-starred GitHub repositories that satisfied several additional criteria. First, to reflect contemporary industry practices, we restricted our search to repositories created within the past 15 years⁵ and excluded repositories that had been inactive for more than one year⁶. Second, to ensure representativeness of real-world Java projects, we excluded templates, archived repositories, and forks by specifying appropriate parameters in our GitHub API queries. Third, we required repositories to contain at least 2000 commits to guarantee substantial change history. We further filtered out repositories whose codebase consisted of less than 95% Java, thereby ensuring the availability of a sufficiently large number of Java methods for extraction.

Following the automated filtering process, we manually reviewed all remaining repositories to verify adherence to the above criteria and to confirm that their descriptions were written in English. Ensuring English descriptions facilitates the identification of bug-fixing commits using English-language keywords and supports subsequent manual qualitative analysis of commit messages. Based on this review, we selected 49 additional projects. Combined with the 49 projects from Chowdhury et al. [13], our final dataset comprises 98 projects in total. The complete list of filtered projects is included in our repository⁷.

3.1.2 Extracting Method-level History. To quantify the bug-proneness of a method, it is necessary to determine how frequently that method has been involved in bug-fixing activities, which in turn requires capturing its complete change history. To obtain this historical information, we employed the CodeShovel tool [24]. CodeShovel has demonstrated high accuracy in identifying method-level change histories and is robust to complex code transformations, including method relocations across files [24]. Its reliability has also been independently evaluated by industry practitioners, who reported similarly strong accuracy [24]. Although a more recent tool, CodeTracker [33], has been shown to achieve higher accuracy than CodeShovel, its substantially slower performance renders it impractical for tracing the histories of large numbers of methods. Moreover, CodeTracker has not undergone independent validation by industry professionals.

³<https://docs.github.com/en/rest>

⁴<https://github.com/PyGithub/PyGithub>

⁵Created after January 1st, 2010

⁶No commits since May 13, 2024

⁷<https://github.com/SQMLab/ExtremelyBuggyPublicData/tree/main/AdditionalFigures/RepositoryList.md>

To support large-scale and reliable analysis, we used CodeShovel to generate structured JSON outputs for approximately 1.25 million methods across 98 projects, capturing their complete method-level evolution across commits. These outputs provide a robust foundation from which we extract both code metrics and change metrics for every method revision included in our dataset.

3.2 Labeling Bug-fix Commits

The bug-proneness of a method is historically estimated by extracting bug-related keywords from commit messages. A commonly used set of such keywords (and their variations) includes: *error*, *bug*, *fix*, *issue*, *mistake*, *incorrect*, *fault*, *defect*, and *flaw* [42, 58]. However, this keyword set often produces false positives, primarily due to the inclusion of the keyword *issue* [13]. To reduce these false positives, Chowdhury *et al.* [13] proposed removing the keyword *issue* and focusing on commit messages that contain both bug-related and fix-related keywords. Although this method significantly improves the accuracy of identifying bug-fix commits, it does not eliminate noise caused by tangled commits—commits that include changes to both bug-related and unrelated methods—leading to incorrectly labeled methods.

While handling tangled code changes remains an active research area [13, 27, 48], we conduct our analysis using three distinct datasets to mitigate the impact of tangled commits. Our underlying hypothesis is that if similar results are observed across all three datasets, the findings are more robust and less likely to be impacted by tangled changes.

- (1) **HighRecall Dataset.** Here, we used all the above-mentioned keywords and their variations (include *issue*) and did not consider the problem of tangled changes. While this approach obtains high recall since any change remotely related to a bug-fix is labeled as a bug-fix, it results in low precision as many changes not related to a bug-fix are labeled as such.
- (2) **HighPrecision Dataset.** In this dataset we used the set of keywords introduced by Chowdhury *et al.* [13], which aims at increasing the precision of bug-fix labeling. To remove the effects of tangled changes, we only labeled a change as a bug-fix if at most 1 method was changed in that commit—if there is only one modified method in a bug-fix commit, that method was certainly a bug-prone method [13]. While this approach achieves a high level of precision, it ultimately reduces recall, as many bug fixes may be ignored.
- (3) **Balanced Dataset.** We used the same set of keywords as the HighPrecision Dataset; however, we allowed up to 5 methods changed in a commit for it to be labeled as a bug fix. This approach aims at achieving a balance between precision and recall.

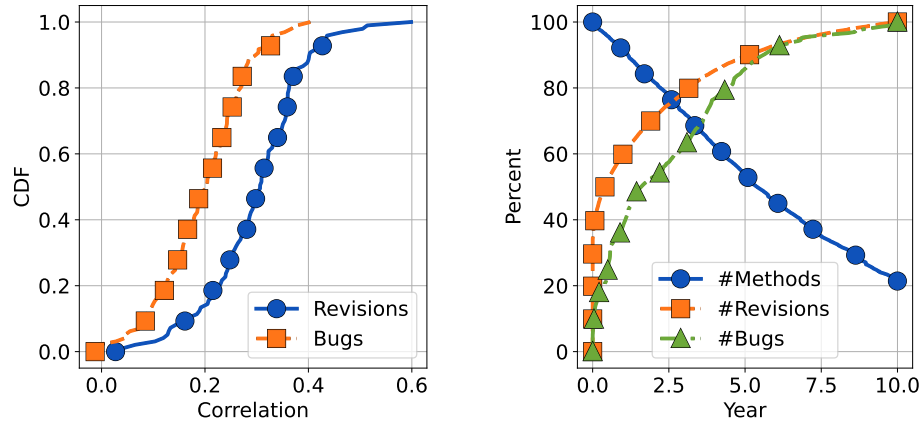
3.3 Age Normalization

The age of the methods in our dataset varies significantly. This means that comparing newer methods to older methods may introduce bias, as older methods have more time to undergo changes and hence, undergo bug-fix related changes. This assumption is confirmed in figure 1a, where we can see that for many projects, there is a positive correlation when comparing age against both revisions and bugs. Specifically, for approximately 60% of projects, the correlation between a method’s age and its number of bugs is ≥ 0.2 —while weak, such correlations still should not be ignored to reduce confounding in our analysis.

Due to this correlation, we removed all methods that are under 5 years old from our dataset. However, we still cannot directly compare methods that are 5 years old with those that are more than 5 years old for the aforementioned reasons. Thus, we also remove all changes made after the first 5 years from the remaining methods. This, unfortunately, resulted

in us losing 626,529 methods and all methods from the following 6 projects from our dataset: 'junit5', 'SmartTube', 'RxJava', 'Essentials', 'titan', 'xpipe'. Since we are interested in the Extremely Bug-Prone methods, we care more about retaining change history—allowing methods to exhibit their true bug-proneness. Thus, the reasoning behind a 5-year threshold can be summarized in Figure 1b. Here we see that 5 years allows us to retain ~90% of changes and ~85% of bugs while still saving ~50% of the methods in our dataset. If we decrease the threshold, we will lose more bug-proneness information in our dataset.

The end result of this age-normalization is a dataset of 691,610 methods from 92 projects, where each method has exactly 5 years of change history. All the projects in our dataset, after age normalization, as well as the number of *ExtremelyBuggy* methods in each project for the Balanced dataset can be viewed in our public repository⁸.



(a) Correlation between age and both revisions and bug-proneness for all 98 projects

(b) Percent of total methods, revisions and bugs retained at different age-normalization thresholds.

Fig. 1. Both figures use the HighRecall Dataset to measure bugs. To enhance readability, we only use 11 markers when presenting each distribution, however, all data points are represented.

3.4 Labeling Bug Proneness

For all three datasets, we assign each method to one of three categories based on the number of bug-fix commits in which it was involved:

- *NotBuggy*: Methods that were never involved in a bug-fix commit.
- *Buggy*: Methods that were involved in exactly one bug-fix commit.
- *ExtremelyBuggy*: Methods that were involved in more than one bug-fix commit.

3.5 Code Metrics

Code metrics have traditionally served as indicators of quality and maintainability of code [13, 21, 50]. Encouraged by previous research in method-level code metrics [12, 41, 50], we used 14 code metrics to determine if the code quality indicators of the *ExtremelyBuggy* methods are different than others. The selected metrics are described as follows.

⁸<https://github.com/SQMLab/ExtremelyBuggyPublicData/tree/main/AdditionalFigures/RepositoryList.md>

Size. Size is one of the simplest metrics, yet it is considered one of the most important with regard to software quality and maintenance [14, 17, 21]. We used the number of source lines of code without comments and blank lines as our size measurement (referred to as SLOCStandard) [12, 13, 36, 57].

Readability. The ability to read and understand existing code is an important part of software maintenance [62]. Code readability can be broadly defined by how easy it is for a developer to understand the structure and flow of source code [4, 5, 34, 52]. Developers can more easily understand and therefore modify readable code with reduced risk of issues. We used two distinct readability metrics in this paper: *Buse et al.* [4] (*Readability*) and *Posnett et al.* [52] (*SimpleReadability*). Both provide a value for each method in the range [0, 1].

Complexity and Testability. We used the McCabe metric [39] to analyze the complexity of methods. This metric represents the number of independent paths in a method; the greater the number of paths, the more challenging it becomes to test the method thoroughly. In addition, we used the number of control variables (*NVAR*), the number of comparisons (*NCOMP*), and the McClure metric [40], which is calculated as the sum of *NVAR* and *NCOMP*. These metrics offer additional insights into method complexity that are not captured by the McCabe metric. Hindle *et al.* [29] proposed the proxy indentation metric (*IndentSTD*), which performs similarly to the McCabe metric in capturing code complexity but, unlike McCabe, does not require a language-specific parser. We included *IndentSTD* in our metric set as well. Finally, we used the *MaximumBlockDepth* metric, which measures the maximum nested block depth, as deep nesting within a method can further complicate testing.

Dependency. When methods are highly coupled, a bug within one method is more likely to affect another. To measure this dependency between methods, we used *totalFanOut*, which is the total number of methods called by a given method.

Maintainability. *Maintainability Index* (MI) [47] is a composite software quality metric that integrates *HalsteadVolume*, which quantifies the amount of information a reader must process to comprehend the code’s functionality, along with *McCabe* complexity and *Size*. The MI is computed using the following formula:

$$MI = 171 - 5.2 \times \ln(HalsteadVolume) - 0.23 \times (McCabe) - 16.2 \times \ln(Size) \quad (1)$$

It produces a single score ($-\infty, 171$], designed to reflect the maintainability of a given method, where a higher score reflects better maintainability. Industry standard tools such as Visual Studio⁹ use this metric, making it a worthy inclusion for our purposes.

Other. Halstead Length (*Length*), which measures the total number of operators and operands, the number of parameters (*Parameters*) and the number of local variables (*LocalVariables*), are also included alongside the other metrics. We also intended to include the *FanIn* metric in our analysis, but were unable to do so due to the complexity involved. For example, consider the Hadoop project, which contains approximately 70,000 methods. Since our goal is to make predictions at method inception, we need to capture code metrics at the time each method was first introduced. Measuring the *FanIn* metric would require constructing a separate call graph for the whole codebase at each method’s introduction commit. Unfortunately, building such call graphs is extremely time-consuming, as it requires preprocessing the entire codebase. Given that we work with 98 projects—each potentially requiring thousands of unique call graphs corresponding to thousands of distinct commits—this approach was computationally infeasible within the scope of our study.

⁹<https://learn.microsoft.com/en-us/visualstudio/code-quality/code-metrics-maintainability-index-range-and-meaning?view=vs-2022>

3.6 Statistical Tests

After evaluating our data with the Anderson-Darling normality test [59], we found that most of the distributions in our data (e.g., the code metrics) do not follow a normal distribution. Therefore, we adopted the non-parametric Wilcoxon rank sum test to assess if there is a statistical difference between two given distributions. To understand the size of those difference, we employed the non-parametric Cliff’s d and categorized the effect sizes following Hess et al. [28]: *Negligible* (N) (< 0.147), *Small* (S) ($0.147 \leq \delta < 0.33$), *Medium* (M) ($0.33 \leq \delta < 0.474$), and *Large* (L) ($\delta \geq 0.474$). Both of these statistical tests do not assume normally distributed distributions and have been widely used in software engineering research [1, 8, 10, 26, 51]. In addition, we used the non-parametric Kendall’s τ correlation coefficient instead of Pearson’s r [10, 21, 31]. Unless otherwise stated, we use a significance threshold of $p \leq 0.05$ when evaluating the statistical significance of correlation coefficients.

4 Approach, Analysis and Results

In this section, we discuss the approach of each RQ along with the findings.

4.1 RQ1: Proportion of *ExtremelyBuggy* methods and their impact

Past research has shown that a high proportion of changes come from a small segment of code [9, 23]. Motivated by this observation, we investigate what proportion of methods are *ExtremelyBuggy* and how many bugs the *ExtremelyBuggy* methods account for in a project.

Table 1. Number of *Buggy* and *ExtremelyBuggy* methods in all 3 datasets.

Dataset	# <i>Buggy</i> (%)	# <i>ExtremelyBuggy</i> (%)
HighRecall	79715 (11.53)	45860 (6.63)
HighPrecision	2704 (0.39)	287 (0.04)
Balanced	8024 (1.16)	1195 (0.17)

Table 1 shows the number of *Buggy* and *ExtremelyBuggy* methods for all three different datasets when we consider all 92 projects aggregated together. Although the high recall dataset contains a large number of *ExtremelyBuggy* methods (6.63%), these numbers decline significantly when we observe the HighPrecision and Balanced datasets. Specifically, only 0.04% (or 287 of the ~700000) methods are *ExtremelyBuggy* in the HighPrecision dataset. Another observation is that the ratio of *Buggy* to *ExtremelyBuggy* methods is smaller in the HighPrecision and balanced datasets. In the HighRecall dataset, approximately 37% of all methods with at least one bug-fix (*Buggy* and *ExtremelyBuggy*) are *ExtremelyBuggy* whereas in the HighPrecision and Balanced datasets, this percentage drops to approximately 10% and 13% respectively.

By viewing the results in figure 2, we can see the percent of methods that are *Buggy* and *ExtremelyBuggy* for individual projects across all datasets. Each plot shows the CDF of the percent of methods that are *Buggy* and *ExtremelyBuggy* on a log scale where each data point is an individual project. For example, in the HighPrecision dataset, in ~80% of projects, $\leq 1\%$ of methods are *Buggy*. One interesting finding is that for the HighPrecision dataset, ~35% of projects have no *ExtremelyBuggy* methods. This will make it very hard to do a project-wise analysis on the *ExtremelyBuggy* methods using this dataset since the results will be skewed by the large number of projects with zero *ExtremelyBuggy* methods.

We now investigate what percent of bugs are captured by the *ExtremelyBuggy* methods alone. We counted the number of bugs in the *ExtremelyBuggy* methods in each project and divided it by the total number of bugs in the project

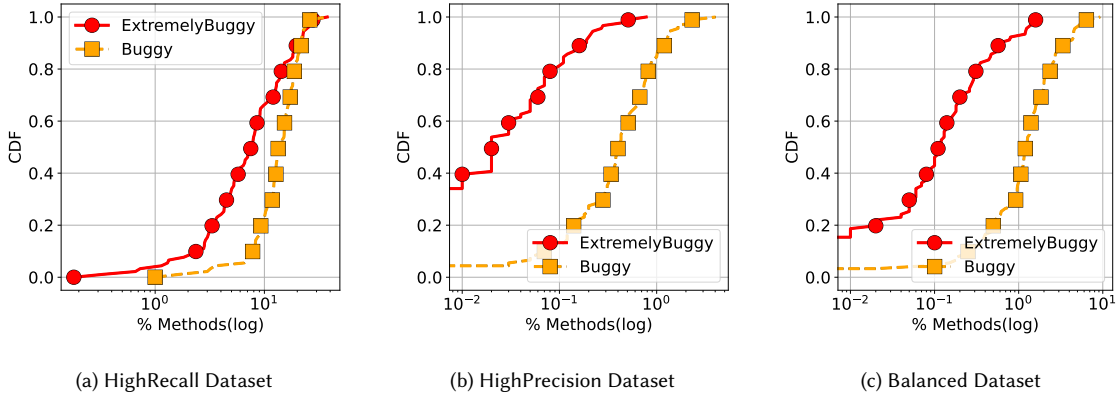


Fig. 2. CDFs of the percent of methods in that are *Buggy* and *ExtremelyBuggy* for each of the 3 datasets for all projects. To enhance readability, we only use 11 markers when presenting each distribution, however, all 92 projects are represented.

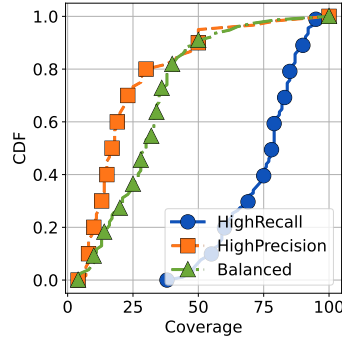


Fig. 3. CDFs of the percent of bugs covered by the *ExtremelyBuggy* methods for each of the three datasets for all 92 projects, excluding projects with no *ExtremelyBuggy* methods. To enhance readability, we only use 11 markers when presenting each distribution.

for each of the 92 projects. To avoid overcounting the number of bugs in each project, we only count a bug once for each unique commit. The results for each dataset are presented in figure 3. For the HighRecall Dataset, we can see that the *ExtremelyBuggy* methods account for a high proportion of bugs, as for ~80% of projects the *ExtremelyBuggy* methods accounted for $\geq 60\%$ of bugs and ~50% of projects the *ExtremelyBuggy* methods covered $\geq 75\%$ of bugs. Not surprisingly, the results are much different for the HighPrecision and Balanced datasets. For these two datasets, we see much less bug coverage of the top percent of methods with bugs. However, this is expected as in the HighPrecision and Balanced datasets, many of the *ExtremelyBuggy* methods were excluded for maintaining higher precision.

RQ1 Summary: Across all datasets, only a small fraction of methods are *ExtremelyBuggy*, yet they account for a disproportionately large share of bugs. In the HighRecall dataset, these rare methods account for a large share of bugs, often more than 60% for 80% of the projects. This underscores the importance of early identification and remediation of these methods, as proactive action is crucial for minimizing future maintenance costs.

4.2 RQ2: Code Quality of *ExtremelyBuggy* Methods

To answer this research question, we analyze how the code quality indicators of *ExtremelyBuggy* methods differ from those of both *Buggy* and *NotBuggy* methods. Since our goal is to identify *ExtremelyBuggy* methods at the time of their introduction, we compare the code metrics of all three categories—*NotBuggy*, *Buggy*, and *ExtremelyBuggy*—at the point when the methods are first added to the codebase.

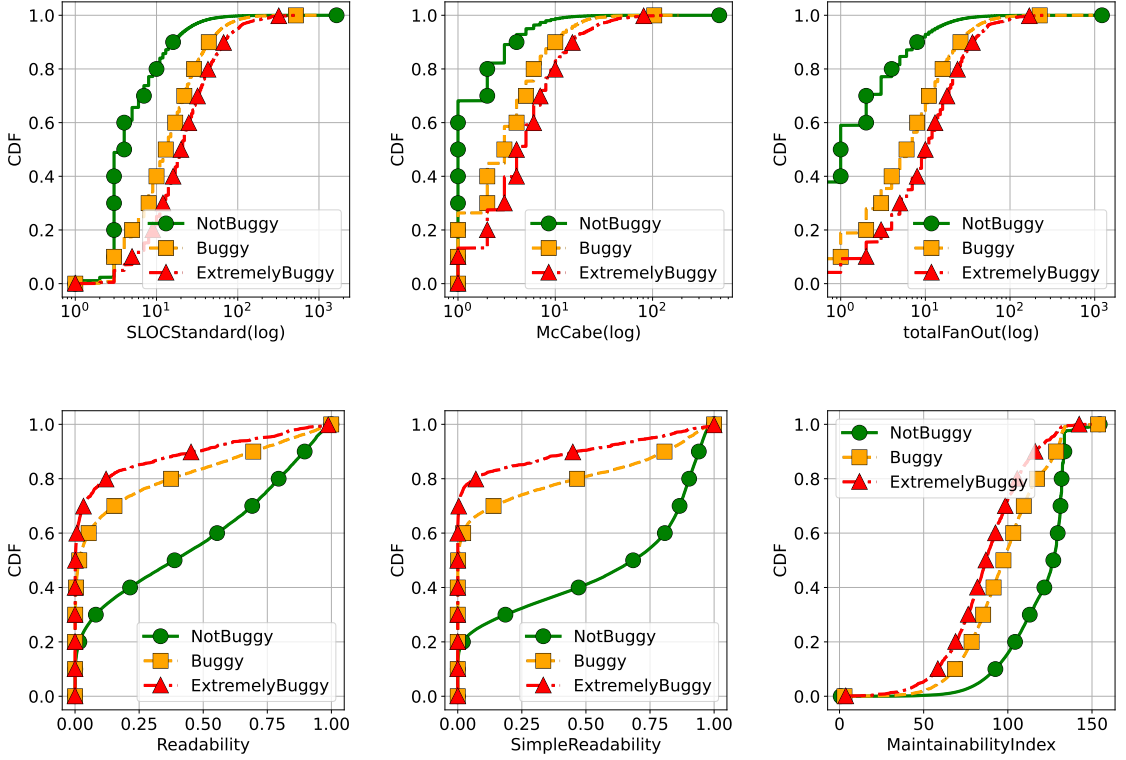


Fig. 4. CDFs showing the distributions of different code metrics at the first commit between *NotBuggy*, *Buggy*, and *ExtremelyBuggy* methods for the Balanced dataset. To enhance readability, we only use 11 markers when presenting each distribution; however, all methods are represented. Note that 34 of the ~700000 methods had a *MaintainabilityIndex* below zero. To improve readability, we excluded these methods when constructing the *MaintainabilityIndex* graph.

Figure 4 compares the three different method classes for six of the selected code metrics using the Balanced dataset; the results are similar for all code metrics across all three datasets. Clearly, *ExtremelyBuggy* methods are larger, less readable, and more complex than both the *Buggy* and *NotBuggy* methods.

To confirm these visual findings, we also performed a statistical analysis to compare *ExtremelyBuggy* methods with both *Buggy* and *NotBuggy* methods. In Table 2 we present the results when comparing the *ExtremelyBuggy* methods to *Buggy* and *NotBuggy* methods grouped together for the Balanced dataset. The comparisons were done on all 92 projects' distributions aggregated together. We can see that all the P-values are significant ($P \leq 0.05$) and that all the effect sizes are *large*, except for *Parameters*, which is *medium*. This indicates that there is a substantial difference

between *ExtremelyBuggy* methods and all other methods at inception. Additionally, the sign of the effect size is negative (–) for *Readability*, *SimpleReadability*, and *MaintainabilityIndex* indicating that these methods are less readable and maintainable. For other code metrics the sign is positive (+) indicating that the *ExtremelyBuggy* methods are larger and more complex. The results for the HighRecall dataset are the same; however, the effect sizes for the HighRecall dataset are mostly *medium*. This is because the HighRecall dataset has a higher number of false positives when identifying the *ExtremelyBuggy* methods, resulting in more noise and less distinction of the *ExtremelyBuggy* methods.

Table 2. The results of statistical significance test (Wilcoxon rank sum test) comparing *ExtremelyBuggy* methods with the combined group of *Buggy* and *NotBuggy* methods for the Balanced dataset dataset. The method groups used in the comparison are formed by aggregating methods from all projects.

Metric	P-value	Sign(+/-)	Effect Size
SLOCStandard	0.00	+	large
Readability	0.00	-	large
SimpleReadability	0.00	-	large
NVAR	0.00	+	large
NCOMP	0.00	+	large
Mcclure	0.00	+	large
McCabe	0.00	+	large
IndentSTD	0.00	+	large
MaximumBlockDepth	0.00	+	large
totalFanOut	0.00	+	large
Length	0.00	+	large
MaintainabilityIndex	0.00	-	large
Parameters	0.00	+	medium
LocalVariables	0.00	+	large

In addition to the aggregated analysis, we also perform individual project analysis to obtain generalizable observations. Projects that has no *ExtremelyBuggy* methods were excluded from this analysis. Table 3 presents the results for the *ExtremelyBuggy* compared to the *Buggy* and *NotBuggy* methods grouped together for the Balanced dataset. For example, for 78.21% of projects, the difference in *SLOCStandard* between *ExtremelyBuggy* and other methods is statistically significant ($P \leq 0.05$). For these projects, the difference was *large* (L) for 98.36% of the projects and for the remaining 1.64% it was *medium* (M). For all code metrics, some projects (usually approximately 20-30%) did not show statistical differences between *ExtremelyBuggy* and other methods. This is because some projects in the Balanced dataset have a very low number of *ExtremelyBuggy* methods, making it difficult to discern differences in distributions. However, for the projects that did have a statistical difference, that difference was almost always *large*. Out of the 14 code metrics, only one, *Parameter*, exhibited distinct characteristics. While more than 62% of projects consistently demonstrated statistical significance for all other metrics, *Parameter* differed for only 20% of projects, indicating that the bug-proneness of methods is not significantly influenced by parameter count.

For the HighPrecision dataset, approximately 40% of projects exhibited statistical differences between *ExtremelyBuggy* and other methods for all metrics, with the effect size remaining almost always large. For the HighRecall dataset, approximately 90-95% of projects demonstrated statistical differences for all metrics; however, the effect size ranged from small (S) to large (L) (typically ~30% in each category). This is expected due to the high number of *ExtremelyBuggy* methods but lower precision in identifying *ExtremelyBuggy* methods present in the HighRecall dataset.

Table 3. The results of statistical significance test (Wilcoxon rank sum test) and effect size (Cliff’s delta) comparing *ExtremelyBuggy* methods with the combined group of *Buggy* and *NotBuggy* methods for the Balanced dataset dataset, evaluated per project. Each value denotes the percentage of projects exhibiting the corresponding behavior. For instance, for the *totalFanOut* metric, 80.77% of projects show a statistically significant difference, and in 98.41% of those cases, the effect size is classified as *large* (L).

Metric	$P \leq 0.05$	N	S	M	L
SLOCStandard	78.21	0.00	0.00	1.64	98.36
Readability	65.38	0.00	3.92	5.88	90.20
SimpleReadability	74.36	0.00	0.00	1.72	98.28
NVAR	66.67	0.00	0.00	7.69	92.31
NCOMP	69.23	0.00	0.00	3.70	96.30
Mcclure	67.95	0.00	0.00	3.77	96.23
McCabe	69.23	0.00	0.00	3.70	96.30
IndentSTD	62.82	0.00	4.08	12.24	83.67
MaximumBlockDepth	69.23	0.00	0.00	5.56	94.44
totalFanOut	80.77	0.00	0.00	3.17	96.83
Length	79.49	0.00	0.00	1.61	98.39
MaintainabilityIndex	80.77	0.00	0.00	1.59	98.41
Parameters	20.51	0.00	25.00	50.00	25.00
LocalVariables	65.38	0.00	1.96	7.84	90.20

To further understand how *ExtremelyBuggy* methods compare to *Buggy* and *NotBuggy* ones, we also did these comparisons without grouping the *Buggy* and *NotBuggy* methods together. We found that the effect size differences between *NotBuggy* and *Buggy* methods were generally *small* to *medium* and the same was found when comparing the *Buggy* to *ExtremelyBuggy* methods. However, the difference between *NotBuggy* and *ExtremelyBuggy* methods was mostly *large* for all code metrics and datasets. This finding confirms the visual findings from figure 4 about the differences between the three bug-proneness types.

RQ2 Summary: At their inception, *ExtremelyBuggy* methods are significantly larger, less readable, and more complex than *Buggy* and *NotBuggy* methods. This is encouraging because it suggests these methods are distinguishable at the very beginning of their lifetime.

4.3 RQ3: Predicting the *ExtremelyBuggy* Methods

In this section, we investigate whether *ExtremelyBuggy* methods can be predicted at the time of their creation. Inspired by prior research [9, 41, 50], we employ five machine learning models of increasing complexity: Logistic Regression [30], Decision Tree [68], Random Forest [60], AdaBoost [63], and a Multi-Layer Perceptron (MLP) model [44]. Our goal is to systematically evaluate the predictive power of both simple and complex models, assessing whether more sophisticated algorithms offer a significant advantage in identifying buggy code early in the development lifecycle. All models were implemented using the default parameters provided by the `scikit-learn` library. This choice was motivated by two factors: (i) hyperparameter tuning is computationally expensive under the leave-one-out validation setting—the evaluation approach used in this paper—as it requires training a separate model for each test instance, and (ii) prior similar studies have found that default settings often yield competitive performance [9, 50]. Moreover, using default parameters enhances the reproducibility of our study and facilitates future replication efforts.

To build and train these models, we first merged the *Buggy* and *NotBuggy* methods into a single category: not *ExtremelyBuggy*. This transformation enables a binary classification task, allowing the models to specifically focus on distinguishing *ExtremelyBuggy* methods from all others. The models were trained using the 14 code metrics computed for each method at the time of its introducing commit. Making predictions by relying solely on metrics available immediately after a method is created allows us to assess the feasibility of detecting *ExtremelyBuggy* methods as early as possible in the software development process.

Due to the highly imbalanced nature of our dataset, we do not use metrics such as *accuracy* to evaluate the performance of the machine learning models, as they could overstate the performance of our models. Instead, we use *precision* and *recall*, *F-measure* (F1), Area Under the Curve (AUC) score, and the Mathews Correlation Coefficient (MCC). AUC ranges from 0 to 1, with a score of 0.5 being no better than random classification and 1 being perfect. MCC ranges from -1 to 1, with 0 being about random classification and a score of 1 being perfect prediction.

We assess prediction performance using a leave-one-out approach at the project level. In each iteration, all methods from one project are used as the test set, while methods from all other projects serve as the training data. This process is repeated for every project, ensuring that each serves as the test set exactly once. This evaluation strategy is motivated by two key factors: (i) it prevents data leakage between training and test sets [13, 50], and (ii) it allows us to assess how well the models generalize across different projects. To address the data imbalance problem—specifically, the significantly smaller number of *ExtremelyBuggy* methods—we apply both undersampling and oversampling techniques using modules provided by the *scikit-learn* library. While the overall prediction performance remains similar across these strategies, undersampling yields slightly better results.

Figure 5 presents various metrics for two selected algorithms using the undersampling approach. The results indicate that recall is substantially higher than precision, suggesting that while the machine learning models are effective at identifying most *ExtremelyBuggy* methods, they tend to generate a high number of false positives. Most project-level evaluations fall within a similar range, but a few notable outliers stand out. For example, at least one project achieved perfect precision (1.00), despite approximately 95% of projects having precision below 0.50. Likewise, although around 90% of projects achieved recall scores above 0.5, there was at least one project with a recall below 0.25.

RQ3 Summary: Predicting *ExtremelyBuggy* methods at inception using simple code metrics and standard machine learning algorithms remains highly challenging. While prediction performance varies across projects, the overall results are not yet reliable enough for practical use.

4.4 RQ4: Observable Characteristics of *ExtremelyBuggy* Methods

The futility of accurate prediction led us to manually characterize the *ExtremelyBuggy* methods to provide insights for both practitioners and the research community. We employed the thematic analysis [18] approach to find themes of the *ExtremelyBuggy* methods. Thematic analysis involves identifying general patterns (or "themes") within the data. These themes are created by labeling the data with codes that capture and characterize individual instances of the data and later identifying themes within the codes. Our thematic analysis approach followed the guidelines of Cruzes *et al* [15], which provides a set of steps for thematic analysis in software engineering research. The approach is described as follows.

¹⁰<https://github.com/SQMLab/ExtremelyBuggyPublicData/tree/main/AdditionalFigures/LeaveOneOutTraining>

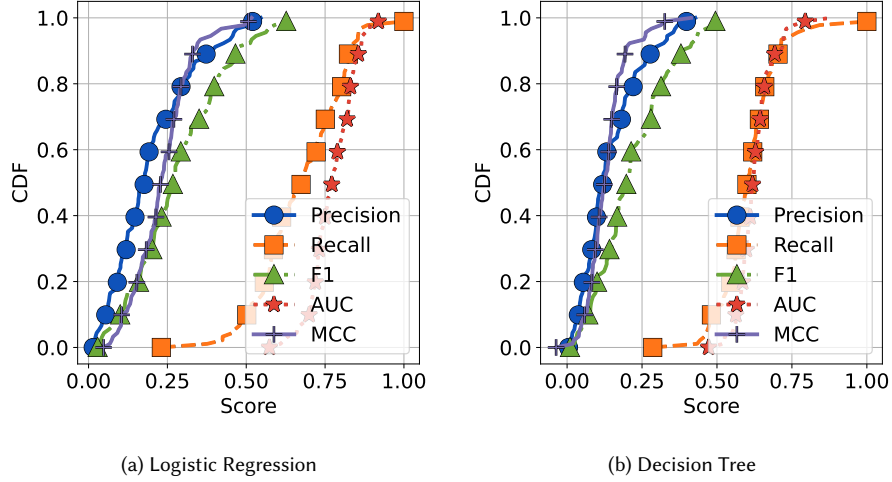


Fig. 5. Results for two of the select machine learning algorithms for the High Recall dataset. Results are very similar with other algorithms, and the results from the other algorithms can be viewed in our publicly shared repository¹⁰. To enhance readability, we use less number of markers than the original number of projects.

Data extraction. For this analysis, we’ve chosen the *ExtremelyBuggy* methods from only the HighPrecision dataset. The reasons behind this are: (i) it contains only 287 unique methods, allowing for thorough manual analysis, and (ii) it has the potential of minimal false positives compared to the other datasets. However, during our review, we still encountered some falsely labeled bugs—mainly due to the inclusion of the keyword *flaw* in commit messages (e.g., “fixing quality flaws”), which sometimes referred to non-functional changes. As a result, we excluded 22 such methods, leaving 265 *ExtremelyBuggy* methods and their evolution history for analysis.

Generating initial codes from the methods. A code is a label that represents some kind of specific, meaningful concept about the data. To help remove internal bias caused by preconceived ideas about the data, we used an open coding [70] approach to generate codes. This approach allowed us not to start with any code and create code intuitively, only creating new code when no existing codes fit the method.

To ensure consistency and correctness in the coding process, the first two authors independently coded 10 selected methods analyzing their history. Following this, the third and fourth authors—who have 2 and 10 years of industry experience, respectively—facilitated a discussion with the first two authors to review the initial codes. We realized that when manually analyzing a data sample, there are multiple dimensions one can take to characterize and assign “codes” to it, as was also discussed in other studies (e.g., [32]). This makes it difficult to generate very similar codes by two independent annotators. For example, in our case, one may choose to look at how a method is written, noting it is very large and hard to read, while another may look at the context it is written in, noting the method is used to parse a piece of text. As such, after the discussion session, we agreed to perform our manual analysis in accordance with the following three distinct dimensions:

- (1) *Visual codes.* These codes correspond to how a method looks. When generating these codes, we asked: *What stands out visually when looking at this method (e.g., is it less readable)?*
- (2) *Context codes.* These codes represent the context a method is used. When generating these codes, we asked: *What is the main purpose of this method (e.g., data handling)?*

- (3) *Bug codes*. These codes are generated to represent the reasons for bugs in a method. When generating these codes, we asked for each bug in a method: *Why did this bug occur (e.g., missing conditionals)?*

The first two authors independently coded an additional set of 10 shared methods and produced highly similar codes. With growing confidence in their consistency, the remaining methods were divided between them for individual coding. While it is difficult to calculate agreement scores in open coding [6], the authors conducted a validity check after analyzing all 265 methods. To assess consistency, each author randomly reviewed 30 methods originally coded by the other, and high agreement was observed—differences appeared in only 12 out of 60 cases. Importantly, these were not cases of conflicting codes; rather, one author had identified additional codes that the other had missed. For instance, one author noted the presence of self-admitted technical debt (SATD) [12, 53], which the other did not. We do not consider these additional codes a limitation. As discussed later, all identified codes, including the extra ones, were collaboratively reviewed by the two authors and synthesized into broader themes, ensuring the robustness and comprehensiveness of the analysis.

Translating our codes into more general themes. Once all our methods had been assigned codes, the first two authors collaborated to identify codes with very similar ideas. They then combined these specific codes into a single, more general theme that accurately encompassed all the codes. It is important to note that a method can belong to multiple themes even within the same theme dimension due to its diverse characteristics.

Creating a model of higher-order themes. Once the more general themes were created, the two authors then worked together to merge themes into higher-order themes that encompass a variety of different codes and themes. These themes were created with the aim of being generic enough to be present in a significant proportion of the identified methods and accurately quantify what the authors observed in the *ExtremelyBuggy* methods, but not too generic that they lose usefulness in characterizing the *ExtremelyBuggy* methods. Some codes and themes had a very small number of methods that contained them, and did not fit into any of the identified higher-order themes. These codes and themes were placed into a higher-order theme labeled as *Other*.

4.4.1 Results. In figure 6, we present an example of the complete thematic analysis process for one of our higher-order themes in the *bug codes* category and the percentages of methods that contain each of the codes and themes. The initial codes that the methods were labeled with are in the right column. In the middle column, we show the initial, more general themes that were identified after looking at the codes. For example, the codes *logging* and *avoiding too many logs* are combined into the more general theme of *logging practices*. Finally, the left column displays the higher-order theme, *exception/error handling and logging* that was identified using the two themes *exception and error handling* and *logging practices*. We combined logging and exception handling into a single broader category because both are commonly used to detect, respond to, and monitor erroneous behavior in program code. Together, they serve as essential tools for maintaining robustness, facilitating debugging, and ensuring system observability during failure scenarios [45].

For each theme, we calculated the *Count* (the number of methods that include this theme), *Percent* (the percentage of all methods that include this theme), and *Mean Percent* (the average percentage of methods containing this theme across all projects). The *Mean Percent* helps determine whether a high overall percentage is driven by a few outlier projects or represents a more generalizable trend.

We now describe the themes and show example methods (CodeShovel generated JSON files) associated with the themes. These example files can be viewed in our public repository.¹¹

¹¹<https://github.com/SQMLab/ExtremelyBuggyPublicData/tree/main/ManualAnalysisExamples>

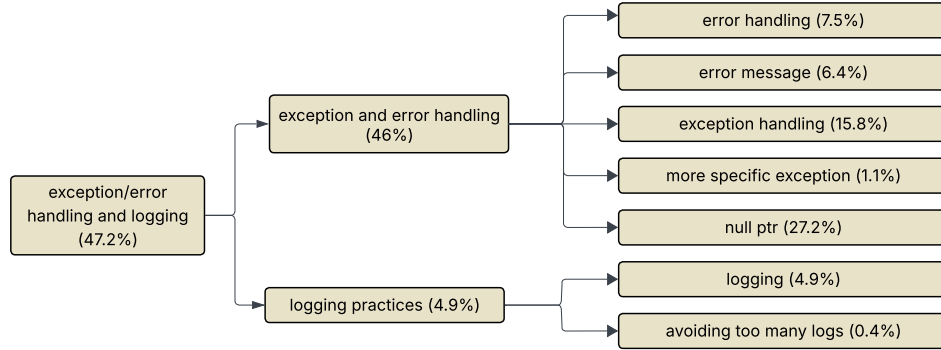


Fig. 6. The taxonomy of one of the themes we identified in the bug themes category during our manual analysis of the *ExtremelyBuggy* methods. The right column shows initial codes, the middle shows more general themes, and the left contains the higher-order theme. Percentages are the number of *ExtremelyBuggy* methods containing that code/theme divided by the total number of *ExtremelyBuggy* methods. The complete taxonomies for all three themes (Bug, Context, and Visual) can be found in the shared repository.

4.4.2 *Visual Themes*. Table 4 shows the final identified themes for the visual/syntactical dimension. These are themes that were identified just by looking at the methods.

Table 4. Final themes from the manual analysis of syntactic and visual characteristics of *ExtremelyBuggy* methods.

Theme	Count	Percent	Mean Percent
Confusing Control Flow	155	58.5	50.8
Abnormal Size	93	35.1	31.7
Other	90	34.0	40.2
Poor Readability	86	32.5	19.5
Drastic Change in Size	80	30.2	37.6
SATD	47	17.7	13.4
Sketchy Exception/Error Handling	35	13.2	14.8

- **Confusing Control Flow.** This theme refers to methods where it is hard to follow the control flow of the method, for example if the method has many if/else statements or deeply nested code. An example is the `processScope()` method (105.json) from the `intellij-community` repository that has many if/else statements and deeply nested code resulting in complicated bugs related to its control flow.
- **Abnormal Size.** These methods are visually very large. This reinforces the common belief that very large methods are substantially more susceptible to bugs than smaller ones. The `processEntry()` method (35.json) from the `jgit` project is an example of a method that is abnormally large.
- **Other.** These are methods that did not show any unique visual characteristics or did not have repeated examples to be grouped into a theme. An example is the `disassemble()` method (166.json) from the `eclipseJdt` project. This means that it may be difficult to detect these methods at inception using source code metrics alone.
- **Poor Readability.** Poor Readability refers to methods that are difficult to read or are formatted very poorly. For example, if a method has bad indentation as seen in the `search()` method (2.json) from the `weka` project.

- **Drastic Change in Size.** These are methods that started with a few lines of code but got much larger over time. These methods are often unimplemented fully at inception and undergo frequent, large changes that result in bugs being introduced later on. An example is the `visitSetOperation()` method (18.json) from the presto project.
- **SATD.** Self-admitted technical debt (SATD) [54] occurs when developers acknowledge that there is future work that needs to be done on a method (e.g., with a TODO comment). An example is the `parseTransition()` method (21.json) from the argouml repository, where unresolved technical debts directly caused bugs in the future.
- **Sketchy Exception/Error Handling.** These methods handled errors and exceptions in dubious ways, often by using many or incomplete try/catch blocks. A notable example is the `invoke()` method (88.json) that contains many try/catch blocks and try blocks without catch blocks.

4.4.3 *Context themes.* Table 5 contains the final 8 themes we identified to explain the contextual purpose of the *ExtremelyBuggy* methods.

Table 5. Final themes from the manual analysis of the contextual purpose of *ExtremelyBuggy* methods.

Theme	Count	Percent	Mean Percent
Core Logic and Algorithms	70	26.4	20.6
Data Handling and Transformation	56	21.1	22.1
Handling External Resources	44	16.6	28.8
Abstract Code Operations	41	15.5	9.6
Error/Exception Handling and Logging	36	13.6	9.8
Initialization and Setup	27	10.2	6.3
UI and Presentation	24	9.1	11.0
Other	18	6.8	7.6

- **Core Logic and Algorithms.** These methods represent some core logical component of a project that determines how that project works. These methods often have complicated logic with requirements that change as the project changes. For example, the method `isInUseableZone()` (249.json), from the mage repository performs specific logic to check if an object is in a zone where it can be used. This method gets updated with more and more logic requirements as the project matures which ultimately result in bugs from incorrect logic.
- **Data Handling and Transformation.** These methods deal with or transform data in some way, for example, by encoding/decoding low-level data or parsing a document. An example is the `scanExpr()` method (178.json) from the xerces2-j repository, which parses an XPath expression and breaks it up into tokens.
- **Handling External Resources.** External resources such as databases and HTTP messages are handled by these methods. External resources often have specific requirements for how messages are sent and received, which can cause issues if not handled correctly. A notable example is the `doPost()` method (161.json) from the tomcat project that has bugs resulting from the incorrect formatting of an HTTP POST message.
- **Abstract Code Operations.** These methods perform operations on abstract code representations, such as building and navigating abstract syntax trees (ASTs). However, most of these methods come from a few projects that deal heavily with ASTs. An example is the `visit()` method (237.json) from the pmd project.
- **Error/Exception Handling and Logging.** The main purpose of these methods is to handle some type of exception/error or to log for debugging purposes. The method `handlePackError()` (80.json) from the jgit repository is an example of this theme. Note that this theme represents the purpose of the method dealing with

error/exception handling, whereas the previous visual theme *Sketchy Exception/Error Handling* describes how the methods performed exception/error handling.

- **Initialization and Setup.** These methods configure and set up objects, which can often result in developers forgetting some parts that they need to configure or initialize fully. An example is the `buildClassifier()` method (181.json) from the weka project, where the developers did not account for missing values during configuration, causing bugs.
- **UI and Presentation.** These methods involve the creation and handling of the visual components in a project. One example is the `createDropDownList()` method (97.json) from the dbeaver repository.

4.4.4 Bug Themes. In this dimension of themes, we focus on what was done to fix a bug, not what the method was doing (context) or how it looks (visual). Table 6 contains the final 10 themes we identified to explain types of bugs in *ExtremelyBuggy* methods.

Table 6. Final themes from the manual analysis of the reasons for bugs in *ExtremelyBuggy* methods.

Theme	Count	Percent	Mean Percent
Conditional Logic	155	58.5	54.5
Exception/Error Handling and Logging	125	47.2	43.3
External Interaction	73	27.5	28.8
Variable Misuse and Naming	54	20.4	19.1
Behavioral Oversights	43	16.2	15.2
Other	28	10.6	11.1
Incorrect Statement Order/Location	25	9.4	2.9
Syntax, Type and Keyword Issues	22	8.3	8.6
Data Lifecycle Mismanagement	21	7.9	6.3
Looping, Indexing and Recursion	20	7.5	4.4

- **Conditional Logic.** These methods had bugs related to incorrect or missing conditional logic. An example is the `update()` method (109.json) from the mindustry repository.
- **Exception/Error Handling and Logging.** These methods generated incorrect error messages, did not properly handle null pointer exceptions, or used the wrong types of exceptions. An example is `getViewSize()` (155.json) from the wicket repository.
- **External Interaction.** These types of bugs occur when a method interacts with other methods and resources. These bugs commonly occur from calling the wrong method, often due to it being overloaded or having a similar name to the one it is trying to call. An example is the method `addDefaultAttributes()` (296.json) from the xerces2-j project that calls `toString()` instead of `stringValue()`.
- **Variable Misuse and Naming.** For these methods, bug-fixes were related to incorrect use and naming of variables. The method `lookupValues()` (44.json) from the pentaho-kettle repository provides an example where the wrong variable was passed to a method.
- **Behavioral Oversights.** This theme represents bugs that were caused by some missed edge case or behavior that was not handled properly. For example, an edge case was missed in the `getHadoopUser()` method (7.json) from the flink repository.
- **Incorrect Statement Order/Location.** This type of bug occurs when the correct statements are present, but are in the wrong order or location in the method. This can often occur due to deep nesting, resulting in a

statement being in the wrong set of brackets. The method `readHeaders()` (145.json) from the `netty` project is one example.

- **Syntax, Type and Keyword Issues.** These bugs occur due to typing issues, usually only requiring one keyword or type to be changed to be fixed. We also noticed that variables were often cast to the incorrect type, which caused the wrong methods to be called or being unable to be called. An example is the method `supresses()` (242.json) from the `pmd` repository.
- **Data Lifecycle Mismanagement.** These are bugs related to the creation, destruction, and handling of data. One example is in the `failed()` (59.json) method from the `pulsar` repository that missed releasing resources, causing a memory leak.
- **Looping, Indexing and Recursion.** These are bugs related to any kind of looping or recursion and indexing problems. For example, in the `replaceEvent()` method (328.json) from the `mage` repo, there is a missing break statement in a loop.

RQ4 Summary: *ExtremelyBuggy* methods often exhibit distinct visual characteristics (e.g., code smells such as large size or poor readability) and have frequent contextual roles (e.g., handling core logic). The underlying causes of bugs—such as missing edge cases or incorrect exception handling—are also commonly shared among these methods. These observations can help practitioners identify areas that warrant greater attention and support researchers in designing more effective features for future machine learning-based bug prediction models.

5 Discussion

Due to the importance of method-level bug prediction, significant research has been done in recent years [13, 19, 20, 41, 50, 65]. However, these earlier studies treated all bug-prone methods as the same, regardless of how many times a method was involved in bug fixes. In this study, we choose to explore the *ExtremelyBuggy* methods as understanding and capturing these methods would enable more optimized resource allocation. We found that *ExtremelyBuggy* methods are very small in numbers but can often account for the majority of bugs in a project (RQ1). This is encouraging because early optimization of these small numbers of methods can significantly reduce future maintenance burdens. Our analysis showed that *ExtremelyBuggy* methods are significantly larger, more complex, and less readable than both *Buggy* and *NotBuggy* methods (RQ2). However, *ExtremelyBuggy* methods proved very hard to predict at their inception, as machine learning models showed poor prediction performance (RQ3). Upon manual analysis, we found that the *ExtremelyBuggy* methods have common visual characteristics, contextual purposes, and similar fix types (RQ4).

5.1 Implications for Researchers

For the research community, we provide a dataset of over 1.25 million Java methods originating from 98 popular open-source projects. To the best of our knowledge, this is the largest dataset on method-level bugs. This enormous dataset not only enables replicating our results, but also can help answer more research questions. While manually analyzing the *ExtremelyBuggy* methods (RQ4), we discovered that over 45% of the bugs in the *ExtremelyBuggy* methods were introduced after modification to the original code. This not only explains the poor prediction performance at the inception of the *ExtremelyBuggy* methods, but also suggests that future research should incorporate the change history of methods while building the prediction models.

Many of the common themes we observed are not detectable from code metrics alone—e.g., HTTP and database operations, SATD, and abstract code operations. This implies that incorporating code embeddings as additional features might help improve the accuracy of these models. We also noticed that prediction performance is not uniform across all projects, as some projects exhibited unusually high or low accuracy. Future work should focus on optimal project selection for training models based on the project under test. Previous work in similar studies has shown the promise of optimal project selection during model training [13, 71].

5.2 Implications for Practitioners

Practitioners seeking optimization opportunities should take into account the contextual and visual themes identified in this paper. For instance, methods associated with a project’s *core logic and algorithms* are significantly more likely to become *ExtremelyBuggy*. When designing solutions for such methods, conducting a focused brainstorming session can help identify the most effective implementation strategies and anticipate potential future requirement changes related to the core logic that might introduce bugs.

Notably, *conditional logic* emerged as one of the most bug-prone elements in methods. Therefore, when developers encounter complex conditional structures—especially in high-risk contexts—they should consider breaking them down into simpler, smaller logical components. Additionally, practitioners are encouraged to use code smell detectors, as many extremely buggy methods exhibit common code smells. Practitioners should also address self-admitted technical debts (SATD) as early as possible. In addition, dedicated effort to the design and implementation of test code can help prevent bugs related to unhandled edge cases and incorrect exception handling that we observed frequently with these methods.

6 Threats to Validity

We have identified multiple multiple threats that may impact the validity of our findings.

Construct validity is impacted by our choice of the CodeShovel tool for extracting method histories. Although CodeShovel has been known to show poor performance for very small methods [12], it has still outperformed other tools, including FinerGit and IntelliJ / git log [24]. A more recent tool, CodeTracker [33], which has been shown to outperform CodeShovel, can be used in future work to see how our results hold up.

We used a variety of popular product metrics for evaluating the *ExtremelyBuggy* methods, however, other metrics could be considered. For example, developer-centric metrics could be used, although we chose not to include developer-centric metrics as they are not very useful at the beginning of a project, and we wanted our findings to be applicable at a projects inception.

To detect the bug-proneness of a method, we used a keyword-based approach, which is known to suffer from false positives [3]. However, this threat was mitigated by the construction of three different datasets: one with high precision, one with high recall, and one aimed at a balance between precision and recall.

To identify themes in the *ExtremelyBuggy* methods, we used a manual labeling approach, which is inherently susceptible to human bias and error. To mitigate this threat, we employed the help of two graduate students, with 2 and 10 years of industry experience respectively. After two of the authors labeled 10 methods, these graduate students analyzed the results, and verified the correctness of the assigned labels. The use of two separate authors for manual labeling of methods may also be a threat to validity if they don’t have agreement between the assignment of labels. To curb this threat, the two authors, labeled 10 of the same methods independently, then discussed together and resolved

any disagreements that arose. They did this again for another 10 methods, and, after finding almost no disagreements, conducted the rest of the analysis separately.

For the manual analysis of the methods, the high precision dataset was used to reduce the number of false positives during labeling. However, the high precision dataset can only consider a commit to be a bug-fix if there was no tangled changes, that is, if only one method was changed in that commit. This may introduce bias towards certain types of bugs that only affect one method. Future work can be done to investigate *ExtremelyBuggy* methods bugs that affect multiple methods to validate our results. Additionally, the *ExtremelyBuggy* methods in the high precision dataset are not evenly distributed, with over half of all *ExtremelyBuggy* methods coming from just 5 projects. However, upon removing these projects we found that none of the main takeaways from the paper changed, in turn mitigating this threat.

Internal validity is hampered due to our selection of statistical tests such as Kendall's τ , the Wilcoxon Rank Sum test and Cliff's d . However, all three are well established and used in software engineering research [11, 12, 21, 31].

External validity is affected due to our use of only open source, java repositories, and it is unclear how our findings would extend to closed source projects or projects in different languages. We also ensured that the descriptions for our repositories were written in english. This was necessary to ensure bug-fix commits were able to be identified using an english language keyword based approach and to ensure we could use commit messages and comments to assist in manual labeling. Future work is needed to find and explore the differences in repositories with commit messages written in different languages.

Conclusion validity of this study is influenced by all previous threats to validity. Additionally, our selection of a five-year threshold for age-normalization was chosen arbitrarily to allow for more change history of methods. Consequently, we repeated most of our experiments using a two and eight year threshold, and none of the major results from the first three RQs changed. However, future work is needed to see how the results may change for the manual analysis of *ExtremelyBuggy* methods using different age-normalization thresholds.

7 Conclusion

This study provides the first large-scale, method-level investigation into *ExtremelyBuggy* methods, those repeatedly involved in bug-fix activities, using a dataset of over 1.25 million methods from 98 open-source Java systems. Our empirical analysis demonstrates that although *ExtremelyBuggy* methods constitute only a small fraction of a project's methods, they frequently account for a disproportionately large share of bugs. These methods exhibit clear distinguishing characteristics at their inception: they are typically larger, more complex, less readable, and less maintainable than both simple buggy and non-buggy methods. Yet, despite these measurable differences, our evaluation of multiple machine learning models reveals that early prediction of *ExtremelyBuggy* methods remains highly challenging due to severe class imbalance, cross-project variability, and the fact that many defects arise through later code evolution rather than initial implementation. These findings highlight limitations in existing metric-based prediction approaches and emphasize the need for richer representations of code and project history.

To complement the quantitative analyses, our manual thematic examination of 265 *ExtremelyBuggy* methods reveals recurring visual, contextual, and defect-related patterns. Visually, these methods often exhibit clear code smells, including confusing control flow, abnormal size growth, poor readability, self-admitted technical debt, and fragile exception-handling structures. Many *ExtremelyBuggy* methods implement core logic, transform complex data structures, or interface with external resources—roles inherently prone to evolving requirements and error-prone interactions. Common root causes of their repeated failures include missing or incorrect conditional logic, improper exception handling, misuse of variables, and interaction errors stemming from related APIs. These observations carry practical

implications: developers may benefit from additional scrutiny and testing of methods that embody these high-risk characteristics, while researchers should explore hybrid models that integrate code embeddings, change-history signals, and contextual semantic information. Overall, this work deepens the understanding of extreme bug-prone methods and establishes a foundation for future research toward more robust, early identification of methods likely to cause recurring bugs.

References

- [1] Abdul Ali Bangash, Hareem Sahar, Abram Hindle, and Karim Ali. 2020. On the time-based conclusion stability of cross-project defect prediction models. *Empirical Software Engineering* 25, 6 (2020), 5047–5083.
- [2] V.R. Basili, L.C. Briand, and W.L. Melo. 1996. A validation of object-oriented design metrics as quality indicators. *IEEE Transactions on Software Engineering* 22, 10 (Oct. 1996), 751–761. doi:10.1109/32.544352
- [3] Christian Bird, Adrian Bachmann, Eirik Aune, John Duffy, Abraham Bernstein, Vladimir Filkov, and Premkumar Devanbu. 2009. Fair and balanced? bias in bug-fix datasets. In *Proceedings of the 7th joint meeting of the European Software Engineering Conference and the ACM SIGSOFT Symposium on the Foundations of Software Engineering*. 121–130.
- [4] Raymond PL Buse and Westley R Weimer. 2009. Learning a metric for code readability. *IEEE Transactions on software engineering* 36, 4 (2009), 546–558.
- [5] Jürgen Börstler and Barbara Paech. 2016. The Role of Method Chains and Comments in Software Readability and Comprehension—An Experiment. *IEEE Transactions on Software Engineering* 42, 9 (Sept. 2016), 886–898. doi:10.1109/TSE.2016.2527791
- [6] M Ariel Cascio, Eunlye Lee, Nicole Vaudrin, and Darcy A Freedman. 2019. A team-based approach to open coding: Considerations for creating intercoder consensus. *Field methods* 31, 2 (2019), 116–130.
- [7] Celerity. [n. d.]. The True Cost of a Software Bug: Part One. <https://www.celerity.com/insights/the-true-cost-of-a-software-bug>. [Online; last accessed 01-Sep-2022].
- [8] Yaohui Chen, Peng Li, Jun Xu, Shengjian Guo, Rundong Zhou, Yulong Zhang, Tao Wei, and Long Lu. 2020. Savior: Towards bug-driven hybrid testing. In *2020 IEEE Symposium on Security and Privacy (SP)*. IEEE, 1580–1596.
- [9] Shaiful Chowdhury. 2025. The Good, the Bad, and the Monstrous: Predicting Highly Change-Prone Source Code Methods at Their Inception. *ACM Transactions on Software Engineering and Methodology* (Jan. 2025), 3715006. doi:10.1145/3715006
- [10] Shaiful Chowdhury, Stephanie Borle, Stephen Romansky, and Abram Hindle. 2019. Greenscaler: training software energy models with automatic test generation. *Empirical Software Engineering* 24 (2019), 1649–1692.
- [11] Shaiful Chowdhury, Reid Holmes, Andy Zaidman, and Rick Kazman. 2022. Revisiting the debate: Are code metrics useful for measuring maintenance effort? *Empirical Software Engineering* 27, 6 (Nov. 2022), 158. doi:10.1007/s10664-022-10193-8
- [12] Shaiful Chowdhury, Hisham Kidwai, and Muhammad Asaduzzaman. 2025. Evidence is All We Need: Do Self-Admitted Technical Debts Impact Method-Level Maintenance?. In *2025 IEEE/ACM 22nd International Conference on Mining Software Repositories (MSR)*. IEEE, 813–825.
- [13] Shaiful Chowdhury, Gias Uddin, Hadi Hemmati, and Reid Holmes. 2024. Method-level Bug Prediction: Problems and Promises. *ACM Transactions on Software Engineering and Methodology* 33, 4 (May 2024), 1–31. doi:10.1145/3640331
- [14] Shaiful Alam Chowdhury, Gias Uddin, and Reid Holmes. 2022. An empirical study on maintainable method size in java. In *Proceedings of the 19th International Conference on Mining Software Repositories*. 252–264.
- [15] D. S. Cruzes and T. Dyba. 2011. Recommended Steps for Thematic Synthesis in Software Engineering. In *2011 International Symposium on Empirical Software Engineering and Measurement*. IEEE Comp Soc, IEEE, LOS ALAMITOS, 275–284.
- [16] Marco D'Ambros, Michele Lanza, and Romain Robbes. 2010. An extensive comparison of bug prediction approaches. In *2010 7th IEEE working conference on mining software repositories (MSR 2010)*. IEEE, 31–41.
- [17] Khaled El Emam, Saïda Benlarbi, Nishith Goel, and Shesh N. Rai. 2001. The confounding effect of class size on the validity of object-oriented metrics. *IEEE Transactions on Software Engineering* 27, 7 (2001), 630–650.
- [18] Jennifer Fereday and Eimear Muir-Cochrane. 2006. Demonstrating rigor using thematic analysis: A hybrid approach of inductive and deductive coding and theme development. *International journal of qualitative methods* 5, 1 (2006), 80–92.
- [19] Rudolf Ferenc, Péter Gyimesi, Gábor Gyimesi, Zoltán Tóth, and Tibor Gyimóthy. 2020. An automatically created novel bug dataset and its validation in bug prediction. *Journal of Systems and Software* 169 (Nov. 2020), 110691. doi:10.1016/j.jss.2020.110691
- [20] Emanuel Giger, Marco D'Ambros, Martin Pinzger, and Harald C. Gall. 2012. Method-level bug prediction. In *Proceedings of the ACM-IEEE international symposium on Empirical software engineering and measurement (ESEM '12)*. Association for Computing Machinery, New York, NY, USA, 171–180. doi:10.1145/2372251.2372285
- [21] Yossi Gil and Gal Lalouche. 2017. On the correlation between size and metric validity. *Empirical Software Engineering* 22, 5 (Oct. 2017), 2585–2611. doi:10.1007/s10664-017-9513-5
- [22] Todd L Graves, Alan F Karr, James S Marron, and Harvey Siy. 2002. Predicting fault incidence using software change history. *IEEE Transactions on software engineering* 26, 7 (2002), 653–661.

- [23] Abraham Grosfeld-Nir, Boaz Ronen, and Nir Kozlovsky. 2007. The Pareto managerial principle: when does it apply? *International Journal of Production Research* 45, 10 (2007), 2317–2325.
- [24] Felix Grund, Shaiful Alam Chowdhury, Nick C. Bradley, Braxton Hall, and Reid Holmes. 2021. CodeShovel: Constructing Method-Level Source Code Histories. In *2021 IEEE/ACM 43rd International Conference on Software Engineering (ICSE)*. IEEE, Madrid, ES, 1510–1522. doi:10.1109/ICSE43902.2021.00135
- [25] Hideaki Hata, Osamu Mizuno, and Tohru Kikuno. 2012. Bug prediction based on fine-grained module histories. In *2012 34th international conference on software engineering (ICSE)*. IEEE, 200–210.
- [26] Peng He, Bing Li, Xiao Liu, Jun Chen, and Yutao Ma. 2015. An empirical study on software defect prediction with a simplified metric set. *Information and Software Technology* 59 (2015), 170–190.
- [27] Steffen Herbold, Alexander Trautsch, Benjamin Ledel, Alireza Aghamohammadi, Taher A Ghaleb, Kuljit Kaur Chahal, Tim Bossenmaier, Bhaveet Nagaria, Philip Makedonski, Matin Nili Ahmadabadi, et al. 2022. A fine-grained data set and analysis of tangling in bug fixing commits. *Empirical Software Engineering* 27, 6 (2022).
- [28] Melinda R Hess and Jeffrey D Kromrey. 2004. Robust confidence intervals for effect sizes: A comparative study of Cohen’sd and Cliff’s delta under non-normality and heterogeneous variances. In *annual meeting of the American Educational Research Association*, Vol. 1.
- [29] Abram Hindle, Michael W Godfrey, and Richard C Holt. 2008. Reading beside the lines: Indentation as a proxy for complexity metric. In *2008 16th IEEE International Conference on Program Comprehension*. IEEE, 133–142.
- [30] David W Hosmer Jr, Stanley Lemeshow, and Rodney X Sturdivant. 2013. *Applied logistic regression*. John Wiley & Sons.
- [31] Laura Inozemtseva and Reid Holmes. 2014. Coverage is not strongly correlated with test suite effectiveness. In *Proceedings of the 36th international conference on software engineering*. 435–445.
- [32] Sigma Jahan, Saurabh Singh Rajput, Tushar Sharma, and Mohammad Masudur Rahman. 2025. Why Attention Fails: A Taxonomy of Faults in Attention-Based Neural Networks. *arXiv preprint arXiv:2508.04925* (2025).
- [33] Mehran Jodavi and Nikolaos Tsantalis. 2022. Accurate method and variable tracking in commit history. In *Proceedings of the 30th ACM Joint European Software Engineering Conference and Symposium on the Foundations of Software Engineering (ESEC/FSE 2022)*. Association for Computing Machinery, New York, NY, USA, 183–195. doi:10.1145/3540250.3549079
- [34] John Johnson, Sergio Lubo, Nishitha Yedla, Jairo Aponte, and Bonita Sharif. 2019. An empirical study assessing source code readability in comprehension. In *2019 IEEE International conference on software maintenance and evolution (ICSME)*. IEEE, 513–523.
- [35] Eirini Kalliamvakou, Georgios Gousios, Kelly Blincoe, Leif Singer, Daniel M German, and Daniela Damian. 2014. The promises and perils of mining github. In *Proceedings of the 11th working conference on mining software repositories*. 92–101.
- [36] Davy Landman, Alexander Serebrenik, and Jorgen Vinju. 2014. Empirical analysis of the relationship between CC and SLOC in a large corpus of Java methods. In *2014 IEEE International Conference on Software Maintenance and Evolution*. IEEE, 221–230.
- [37] Michele Lanza, Andrea Mocci, and Luca Ponzanelli. 2016. The tragedy of defect prediction, prince of empirical software engineering research. *IEEE Software* 33, 6 (2016), 102–105.
- [38] Valentina Lenarduzzi, Alberto Sillitti, and Davide Taibi. 2017. Analyzing forty years of software maintenance models. In *2017 IEEE/ACM 39th international conference on software engineering companion (ICSE-C)*. IEEE, 146–148.
- [39] Thomas J McCabe. 1976. A complexity measure. *IEEE Transactions on software Engineering* 4 (1976), 308–320.
- [40] Carma L McClure. 1978. A model for program complexity analysis. In *Proceedings of the 3rd international conference on Software engineering*. 149–157.
- [41] Ran Mo, Shaozhi Wei, Qiong Feng, and Zengyang Li. 2022. An exploratory study of bug prediction at the method level. *Information and software technology* 144 (2022), 106794.
- [42] Mockus and Votta. 2000. Identifying reasons for software changes using historic databases. In *Proceedings International Conference on Software Maintenance ICSM-94*. IEEE Comput. Soc. Press, Victoria, BC, Canada, 120–130. doi:10.1109/ICSM.2000.883028
- [43] Manishankar Mondal, Banani Roy, Chanchal K. Roy, and Kevin A. Schneider. 2019. Investigating the relationship between evolutionary coupling and software bug-proneness. In *Proceedings of the 29th Annual International Conference on Computer Science and Software Engineering (CASCON '19)*. IBM Corp., USA, 173–182.
- [44] Fionn Murtagh. 1991. Multilayer perceptrons for classification and regression. *Neurocomputing* 2, 5-6 (1991), 183–197.
- [45] Suman Nakshatri, Maithri Hegde, and Sahithi Thandra. 2016. Analysis of exception handling patterns in Java projects: An empirical study. In *Proceedings of the 13th International Conference on Mining Software Repositories*. 500–503.
- [46] Alberto S Nuñez-Varela, Héctor G Pérez-Gonzalez, Francisco E Martínez-Perez, and Carlos Soubervielle-Montalvo. 2017. Source code metrics: A systematic mapping study. *Journal of Systems and Software* 128 (2017), 164–197.
- [47] Paul Oman and Jack Hagemester. 1992. Metrics for assessing a software system’s maintainability. In *Proceedings Conference on Software Maintenance 1992*. IEEE Computer Society, 337–338.
- [48] Md Nahidul Islam Opu, Shaowei Wang, and Shaiful Chowdhury. 2025. LLM-Based Detection of Tangled Code Changes for Higher-Quality Method-Level Bug Datasets. *arXiv preprint arXiv:2505.08263* (2025).
- [49] Fabio Palomba, Andy Zaidman, Rocco Oliveto, and Andrea De Lucia. 2017. An Exploratory Study on the Relationship between Changes and Refactoring. In *Proceedings of the 25th International Conference on Program Comprehension* (Buenos Aires, Argentina). 176–185.

- [50] Luca Pascarella, Fabio Palomba, and Alberto Bacchelli. 2020. On the performance of method-level bug prediction: A negative result. *Journal of Systems and Software* 161 (March 2020), 110493. doi:10.1016/j.jss.2019.110493
- [51] Fabiano Pecorelli, Gemma Catolino, Filomena Ferrucci, Andrea De Lucia, and Fabio Palomba. 2020. Testing of mobile applications in the wild: A large-scale empirical study on android apps. In *Proceedings of the 28th international conference on program comprehension*. 296–307.
- [52] Daryl Posnett, Abram Hindle, and Premkumar Devanbu. 2011. A simpler model of software readability. In *Proceedings of the 8th working conference on mining software repositories*. 73–82.
- [53] A. Potdar and E. Shihab. 2014. An exploratory study on self-admitted technical debt. In *Software Maintenance and Evolution (ICSME), 2014 IEEE International Conference on*. IEEE, 91–100.
- [54] Aniket Potdar and Emad Shihab. 2014. An exploratory study on self-admitted technical debt. In *2014 IEEE International Conference on Software Maintenance and Evolution*. IEEE, 91–100.
- [55] Foyzur Rahman and Premkumar Devanbu. 2013. How, and why, process metrics are better. In *2013 35th international conference on software engineering (ICSE)*. IEEE, 432–441.
- [56] Md Saidur Rahman and Chanchal K Roy. 2017. On the relationships between stability and bug-proneness of code clones: An empirical study. In *2017 IEEE 17th International Working Conference on Source Code Analysis and Manipulation (SCAM)*. IEEE, 131–140.
- [57] Paul Ralph and Ewan Tempero. 2018. Construct validity in software engineering research and software metrics. In *Proceedings of the 22nd International Conference on Evaluation and Assessment in Software Engineering 2018*. 13–23.
- [58] Baishakhi Ray, Vincent Hellendoorn, Saheel Godhane, Zhaopeng Tu, Alberto Bacchelli, and Premkumar Devanbu. 2016. On the "naturalness" of buggy code. In *Proceedings of the 38th International Conference on Software Engineering (ICSE '16)*. Association for Computing Machinery, New York, NY, USA, 428–439. doi:10.1145/2884781.2884848
- [59] Normadiah Mohd Razali, Yap Bee Wah, et al. 2011. Power comparisons of shapiro-wilk, kolmogorov-smirnov, lilliefors and anderson-darling tests. *Journal of statistical modeling and analytics* 2, 1 (2011), 21–33.
- [60] Steven J Rigatti. 2017. Random forest. *Journal of insurance medicine* 47, 1 (2017), 31–39.
- [61] Daniele Romano and Martin Pinzger. 2011. Using source code metrics to predict change-prone Java interfaces. In *2011 27th IEEE International Conference on Software Maintenance (ICSM)*. 303–312. doi:10.1109/ICSM.2011.6080797 ISSN: 1063-6773.
- [62] Simone Scalabrino, Mario Linares-Vasquez, Denys Poshyvanyk, and Rocco Oliveto. 2016. Improving code readability models with textual features. In *2016 IEEE 24th International Conference on Program Comprehension (ICPC)*. IEEE, 1–10.
- [63] Robert E Schapire. 2013. Explaining adaboost. In *Empirical inference: festschrift in honor of vladimir N. Vapnik*. 37–52.
- [64] Emad Shihab, Ahmed E. Hassan, Bram Adams, and Zhen Ming Jiang. 2012. An industrial study on the risk of software changes. In *Proceedings of the ACM SIGSOFT 20th International Symposium on the Foundations of Software Engineering (FSE '12)*. Association for Computing Machinery, New York, NY, USA, 1–11. doi:10.1145/2393596.2393670
- [65] Thomas Shippey, Tracy Hall, Steve Counsell, and David Bowes. 2016. So You Need More Method Level Datasets for Your Software Defect Prediction? Voilà! (ESEM '16).
- [66] Shivkumar Shivaji, E James Whitehead, Ram Akella, and Sunghun Kim. 2012. Reducing features to improve code change-based bug prediction. *IEEE Transactions on Software Engineering* 39, 4 (2012), 552–569.
- [67] Dag IK Sjøberg, Aiko Yamashita, Bente CD Anda, Audris Mockus, and Tore Dybå. 2012. Quantifying the effect of code smells on maintenance effort. *IEEE Transactions on Software Engineering* 39, 8 (2012), 1144–1156.
- [68] Yan-Yan Song and Ying Lu. 2015. Decision tree methods: applications for classification and prediction. *Shanghai archives of psychiatry* (2015).
- [69] Davide Spadini, Fabio Palomba, Andy Zaidman, Magiel Bruntink, and Alberto Bacchelli. 2018. On the Relation of Test Smells to Software Code Quality. In *2018 IEEE International Conference on Software Maintenance and Evolution (ICSME)*. IEEE, Madrid, 1–12. doi:10.1109/ICSME.2018.00010
- [70] Anselm Strauss and Juliet Corbin. 1998. Basics of qualitative research techniques. (1998).
- [71] Zhongbin Sun, Junqi Li, Heli Sun, and Liang He. 2021. CFPS: Collaborative filtering based source projects selection for cross-project defect prediction. *Applied Soft Computing* 99 (2021), 106940.
- [72] Chakkrit Tantithamthavorn and Ahmed E Hassan. 2018. An experience report on defect modelling in practice: Pitfalls and challenges. In *Proceedings of the 40th International conference on software engineering: Software engineering in practice*. 286–295.
- [73] Zixu Wang, Weiyuan Tong, Peng Li, Guixin Ye, Hao Chen, Xiaoqing Gong, and Zhanyong Tang. 2023. BugPre: an intelligent software version-to-version bug prediction system using graph convolutional neural networks. *Complex & Intelligent Systems* 9, 4 (2023), 3835–3855.
- [74] Supatsara Wattanakriengkrai, Patanamon Thongtanunam, Chakkrit Tantithamthavorn, Hideaki Hata, and Kenichi Matsumoto. 2020. Predicting defective lines using a model-agnostic technique. *IEEE Transactions on Software Engineering* 48, 5 (2020), 1480–1496.
- [75] Yuming Zhou, Baowen Xu, and Hareton Leung. 2010. On the ability of complexity metrics to predict fault-prone classes in object-oriented systems. *Journal of Systems and Software* 83, 4 (April 2010), 660–674. doi:10.1016/j.jss.2009.11.704