

How Should We Build A Benchmark?

Revisiting 274 Code-Related Benchmarks For LLMs

Anonymous ACL submission

Abstract

Various benchmarks have been proposed to assess the performance of large language models (LLMs) in different coding scenarios. We refer to them as code-related benchmarks. However, there are no systematic guidelines by which such a benchmark should be developed to assure its quality, reliability, and reproducibility. We propose **How2Bench** comprising a 55-criteria checklist as a set of guidelines to comprehensively govern the development of code-related benchmarks. Using HOW2BENCH, we profiled 274 code-related benchmarks released within the past decade and found concerning issues. Nearly 70% of the benchmarks did not take measures for data quality assurance; over 10% did not even open source or only partially open source. Many highly cited benchmarks have loopholes, including duplicated samples, incorrect reference codes/tests/prompts, and unremoved sensitive/confidential information. Finally, we conducted a human study involving 49 participants and **revealed significant gaps in awareness** of the importance of data quality, reproducibility, and transparency. For ease of use, we provide a **printable version** of HOW2BENCH in Appendix E.

1 Introduction

“If you cannot measure it, you cannot improve it.” — Lord Kelvin (1824-1907)

Recent large language models (LLMs) have shown remarkable capabilities across various domains such as software development (Chen et al., 2021a), question answering (Rogers et al., 2023), and math reasoning (Imani et al., 2023). Various **benchmarks** (Chen et al., 2021a; Jimenez et al., 2024; Austin et al., 2021; Yue et al., 2024; Du et al., 2023a) are proposed to evaluate LLMs’ effectiveness and limitations from multiple perspectives in different application scenarios.

However, **doubts regarding the quality, reliability, and transparency** of various code-related benchmarks arise. For example, a recent study pointed out that “current programming benchmarks are inadequate for assessing the actual correctness of LLM-generated code” (Liu et al., 2023a). Other accusations, including **irreproducible** (Reuel et al., 2024), **closed** data sources (Cao et al., 2024b), **low quality** (Qiu et al., 2024a; Yadav et al., 2024a), and **inadequate validation measures** (Liu et al., 2023a), were also raised, undermining the credibility of these benchmarks and thereby their subsequent evaluation results. This motivates the need for rigorous and thorough **guidelines** to govern code-related benchmark development.

In this paper, we introduce **HOW2BENCH**, a **comprehensive guideline** consisting of a 55-criteria checklist specially designed for code-related benchmarks. This checklist **covers the entire lifecycle** of benchmark development, from *design* and *construction* to *evaluation*, *analysis*, and *release* as shown in Figure 1. It underwent multiple iterations — we initiated a draft inspired by open-source software guidelines (Fogel, 2005) and classical measurement theory (Suppes et al., 1962). We refined it through iterative discussions with practitioners, leading to the finalization of these criteria. HOW2BENCH emphasizes **reliability**, **validity**, **open access**, and **reproducibility** in the benchmark development, ensuring high standards and fostering a more reliable and transparent environment.

Following HOW2BENCH, we conducted an in-depth profiling of 270+ code-related benchmarks developed over the past decade (2014 - 2024). The extent of criteria violations by the profiled benchmarks is concerning:

- Almost 70% of the benchmarks did not take any measures for **data quality assurance**;
- Over 90% did not consider **code coverage** when use passing test cases as an oracle;
- Over half of the benchmarks did not provide

the essential information (e.g., experiment setup, prompts) for **reproducibility**;

- Over 10% are **not open source** or only partially open source.

We observed that even **highly cited benchmarks have loopholes**, including duplicated samples, incorrect reference/tests, unclear displays, and unre-moved sensitive/confidential information. We also observed these loopholes can **propagate**. Over 18% of the benchmarks serve as data sources for subsequent benchmarks (Figure 8). Therefore, the data quality of benchmarks affects their credibility and likely impacts future benchmarks.

To understand the usefulness of the criteria in How2Bench, we conducted a human study involving 49 participants through questionnaires. All participants **concurred on the necessity** of having a checklist for benchmark construction to enhance quality. Nearly all participants with experience in benchmark development **acknowledged the importance of all these 55 criteria**. The study also exposed **gaps in quality awareness**: 16% of participants were unaware of the necessity for data denoising, and over 40% were not aware that the experimental setup and environment could impact the reproducibility and transparency. This paper makes contributions in five aspects:

- **Novelty**. We introduce HOW2BENCH, a comprehensive set of guidelines packaged as a 55-criteria checklist that covers the lifecycle of code-related benchmark development.
- **Significance**. HOW2BENCH presents the first comprehensive set of actionable guidelines for developing high-quality benchmarks, striving to create a more reliable and transparent environment. The human study also highlighted the demand for such a detailed guideline.
- **Usefulness**. HOW2BENCH serves as a guideline for practitioners before/during developing code-related benchmarks, and a checklist for evaluating existing benchmarks after their release. For ease of use, we also provide a **printable version** of HOW2BENCH on Appendix E.
- **Generalizability**. Most criteria listed in HOW2BENCH can be adopted or adapted to other benchmarks such as Question-answering, mathematical reasoning, and multi-modal benchmarks.
- **Long-term Impact**. Our statistics alert the community to the severity and prevalence of non-

standard practices in benchmark development. It ultimately improves the overall quality of benchmarks due to the propagation effect among them.

2 Background

2.1 Code-related Benchmarks

Benchmarks for coding tasks like code generation (Chen et al., 2021a; Austin et al., 2021), defect detection (Just et al., 2014; Gao et al., 2023b; Liu et al., 2024c), and program repair (Jimenez et al., 2024; Risse and Böhme, 2024) are increasingly common, reflecting the growing needs for using LLMs for coding tasks. Recent studies have highlighted various issues with these benchmarks, ranging from design inconsistencies to scope and applicability limitations. For example, (Liu et al., 2023a) found that even some widely used benchmarks, such as HumanEval (Chen et al., 2021a) and MBPP (Austin et al., 2021), contains a non-trivial proportion of bugs in implementation, documentation, and test cases. Our work, in comparison, introduces a detailed guideline that **guides the benchmark development** during the entire lifecycle.

2.2 Related Studies and Surveys

Several recent surveys and empirical studies have profiled the status quo of LLM development. These studies either explore the overall performance for certain areas such as software engineering (Hou et al., 2023; Wang et al., 2024a) or investigate the capabilities of LLMs on specific tasks such as code generation (Dou et al., 2024; Yu et al., 2024) and test generation (Schäfer et al., 2024; Yuan et al., 2024b, 2023a). A survey (Chang et al., 2024) about how to evaluate LLMs was proposed to answer what/where/how to evaluate LLMs. This paper differs from these studies in its purpose and perspectives. Unlike these benchmarks, our work proposed guidelines for future benchmark development and provided a checklist to assess the quality of these existing benchmarks.

Recently, BetterBench (Reuel et al., 2024) is a concurrent work assessing the AI benchmarks against 46 criteria. Then, it scored 24 AI benchmarks in various domains and ranked them. BetterBench differs from this paper in several key aspects: scope (general benchmarks vs. code-related benchmarks), lifecycle division (it addresses benchmark retirement, while How2Bench focuses on benchmark evaluation, analysis, and release), and objectives (scoring benchmarks vs. offering comprehensive guidelines for future benchmark development).

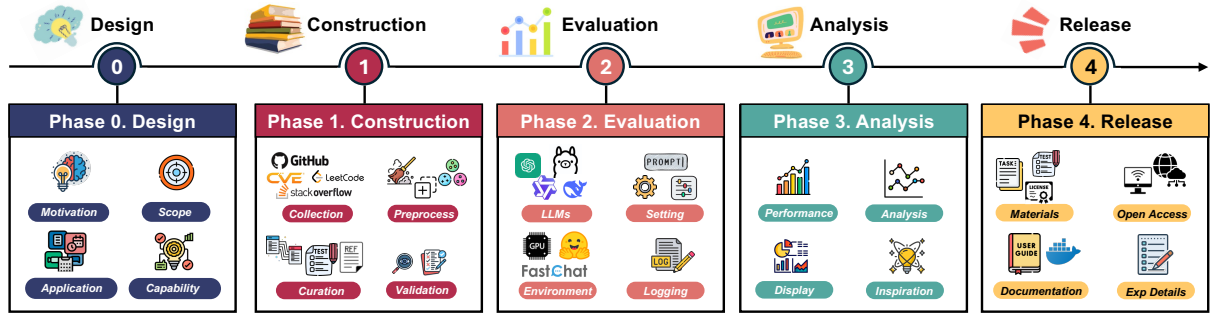


Figure 1: Lifecycle of Benchmark Development

Additionally, the study in this paper was conducted on a much larger scale (24 vs. 274 benchmarks), statistically highlighting the prevalent issues in existing benchmarks.

3 Design

3.1 The Lifecycle of Benchmark Development

Code-related benchmark development comprises five typical phases (Phase 0 - 4), as shown in Figure 1, explained in detail as follows.

Phase 0. Design. At the beginning of benchmark development, it is vital to identify the motivation, the *scope* and the *capabilities* required by the *application* scenario of interest. To achieve this objective, one needs to carefully consider the application scenarios, making sure these scenarios align with real-world demands. Also, it is also necessary to assess whether other benchmarks already exist that address similar tasks, and to identify any shortcomings they may possess. Furthermore, this new benchmark should be designed to evaluate specific LLMs’ capabilities; the crafted tasks are expected to reflect these capabilities.

Phase 1. Construction. After establishing the motivation and purpose, the Benchmark Construction phase moves from design to execution. Typically, data is *collected* from public coding websites such as GitHub, LeetCode and StackOverflow. This is followed by preprocessing, which includes filtering, cleaning (e.g., deduplication, denoising), and *curation* (e.g., aligning tests with corresponding code). The phase usually ends with a *validation* process, which can be manual or automated.

Phase 2. Evaluation. Once the benchmark is available, the next step is to apply it to LLMs, validating if it can effectively measure the intended LLM capabilities. Essential considering factors include *selecting a representative array of LLMs*, configuring *settings* like prompts and hyperparam-

eters for consistency, choosing appropriate experimental *environments* to meet LLM requirements, and implementing thorough *logging* to ensure dependable and reproducible results.

Phase 3. Analysis. After evaluation, experimental results are analyzed, drawing conclusions on LLMs’ capabilities. This phase involves comparing each LLM’s *performance* to identify standout or underperforming models. Then, proper visual aids such as bar charts and tables can be used to *display* the experimental results, presenting clearer observation and deeper *inspiration*, such as the correlations between models, the correlations with related benchmarks, or performance in upper-/downstream tasks. Indeed, a thorough analysis helps pinpoint areas for improvement and guides future enhancements in LLM development.

Phase 4. Release. The final phase is to make the benchmark open-accessible. This phase involves meticulously preparing all *materials* associated with the benchmark, ensuring they are ready for *open access* to foster widespread adoption and collaboration. Clear, comprehensive *documentation* is provided to guide users on effectively utilizing the benchmark. Additionally, all logged *experiment details* are made available, enhancing the reproducibility and transparency of the benchmark.

3.2 Study Design

Our study consists of four steps (Figure 2). All steps are explained as follows.

Step 1. Guideline Construction. To begin with, we *sketched the initial guidelines* for each phase in the benchmark development lifecycle (Section 3.1, Figure 1) by reviewing existing literature (Suppes et al., 1962; Zheng et al., 2023b; Schäfer et al., 2024; Reuel et al., 2024) and brainstorming. After that, we *refined the guidelines* through a series of interviews with various stakeholders, including

or automated testing, and the relevance of these capabilities in real-world *applications*.

■ **Key Statistics** – According to our statistics among 270+ benchmarks, *apparent research bias* can be observed in terms of coding tasks, programming languages, and code granularities are observed (Appendix A.1). For example, 36.13% (99/274) are code generation benchmarks, followed by program repair, with 9.85% (27/274).

Also, during the focused case study (listed in Appendix C), we identified that 10% benchmarks have not explicitly specified the capabilities (e.g., intention understanding, program synthesis) to be evaluated, and 30% *have not specified application scenarios* the benchmark targets.

Besides, we also identified a case in MBPP (Austin et al., 2021) where a case fell out of the target evaluation capabilities (Appendix A.2). Indeed, clearly defining the application scenarios/scopes/capabilities could help benchmark constructors establish precise goals for the design and development of the benchmark, ensuring accuracy in the evaluation.

Lastly, Figure 12 shows that 58% (158/274) code-related benchmarks involve Python, followed by 39% (107/274) involving Java. Yet, 31 programming languages are only covered by one benchmark, and less than five benchmarks cover other 19 programming languages. This observation consolidates the observation from previous works (Cao et al., 2024a; Hou et al., 2023) on a larger scale.

▲ **Severity** – Current benchmarks exhibit *an apparent imbalance* in coding tasks and programming languages dominated by code generation and Python, leaving research blanks to be filled. Also, even highly cited benchmarks may have samples that do not fall into the examined capabilities.

4.2 Guideline for Construction

📖 **Explanation** – Figure 4 shows 19 criteria for benchmark construction. Essentially, for **data source**, the key considerations include verifying the traceability and quality of the data source, addressing potential *data contamination* (Sainz et al., 2023), and ensuring that the *data sampling processes* are scientifically robust and rigorous. Also, for **data representativeness**, it also guides through specific checks to ensure the benchmark’s scope is strictly adhered to, such as making sure every

Phase 1. Benchmark Construction	
5	Consider whether the data source of the benchmark is traceable .
6	Consider whether the data source of the benchmark is of high quality (e.g., stars, downloads, last update times, number of forks).
7	Consider whether the benchmark’s data source is representative (e.g., choose an open-source community or code hosting platform that matches the task, capability, and scope under study).
8	Consider data contamination issues during the benchmark collection (e.g., considering the upload time of the source code or checking whether the data source is included in the training data of LLMs).
9	If data sampling is needed, consider whether the choice of sample size is scientific (e.g., considering the confidence level/margin of error/sampling proportion, etc.).
10	If data sampling is needed, consider whether the sampling process is rigorous (e.g., random sampling, stratified sampling, etc.).
11	Ensure each data point in the benchmark falls into the targeted scope (e.g., checking each data point’s evaluated capabilities or domain knowledge).
12	Consider whether the data in the benchmark can cover the studied capabilities/domain knowledge/application scenarios .
13	Consider whether there is a standard answer for each sample in the benchmark (such as reference code, etc.).
14	For code , consider whether the code is compilable/executable .
15	Consider the possibility of noise in the data and perform denoise .
16	Consider the possibility of duplication in the data and deduplicate them.
17	Clean the sensitive information (such as data desensitization and anonymization) unless the benchmark is deliberately designed so .
18	Manually review some or all of the data in the benchmark to ensure its quality.
19	Use LLMs to review some or all of the data in the benchmark to ensure its quality.
20	Design appropriate output validation methods for the benchmark (e.g., using exact matching or designing test cases).
21	Design appropriate evaluation metrics for the evaluation set (e.g., precision, accuracy, pass@K, recall).
22	Consider the adequacy of the evaluation metrics (e.g., is the code coverage high enough).
23	Consider if there are any other evaluation perspectives (e.g., readability, efficiency, safety, security).

Figure 4: Guideline for Benchmark Construction

data point falls within the targeted scope and that the data can cover all studied capabilities, domain knowledge, and application scenarios.

For **data preprocess and cleaning**, it also stresses handling code-specific aspects, such as compilability and execution, along with cleaning and *manually reviewing* data for quality assurance. Output validation methods and evaluation metrics must be carefully designed and reviewed to ensure they effectively measure the benchmark’s goals. Lastly, it suggests considering additional evaluation perspectives, such as safety (Wei et al., 2024; Yuan et al., 2024a) checks, ensuring the code does not contain sensitive information.

■ **Key Statistics** – According to our statistics (Appendix A.3), the 270+ benchmarks exhibit *numerous irregularities* in their implementation, which could significantly threaten the reliability of the benchmarks. Surprisingly, **62% of benchmarks did not deduplicate** or did not mention. **Near 80% benchmarks did not consider or handle data contamination threats**. **About 70% of the benchmarks did not go through any quality assurance checks** such as manual checks and code execution. In particular, we summarized the *commonly-used data quality assurance metrics* and their frequency: manual check (22.6%), code execution

(2.2%), LLM check (1.5%), others (e.g. the number of stars or heuristic rules, 5.8%).

Also, since we focus on code-related benchmarks, which usually accompany test cases, *test coverage* also needs to be considered. As pointed out by prior study (Liu et al., 2023a), inadequate test coverage can lead to inflated evaluation results. However, we observed that only 8.7% of benchmarks have considered test coverage when using test cases as oracles (Appendix A.3). It severely affects the reliability of findings on these benchmarks, potentially misleading future research and applications based on these flawed assessments.

▲ **Severity** – Most benchmarks display *severe loopholes* in data preparation and curation. Quality checks are often neglected.

4.3 Guideline for Evaluation

Phase 2. Benchmark Evaluation	
24	Consider whether <u>sufficient</u> LLMs are evaluated.
25	Consider whether <u>representative</u> LLMs (e.g., covering latest/classical LLM families, small/large LLMs, and open-/closed-source LLMs) are evaluated.
26	Consider whether the prompt is of <u>high quality</u> (e.g., the instruction and intent are clear).
27	The prompts have been <u>validated by humans or LLMs</u> (e.g., evaluated or discussed by participants or preliminarily tried out on several LLMs).
28	Try <u>different paraphrases</u> of the prompt.
29	Try <u>different prompting strategies</u> to observe the impact on the evaluation results (e.g., in-context learning, chain-of-thought).
30	Pay attention to the <u>hardware environment</u> (such as GPU card, storage size, etc.) of the experiment.
31	Pay attention to the <u>operating system and software environment</u> (e.g. operating system, version, etc.) <u>used for the experiment</u> .
32	Pay attention to the off-the-shelf <u>platforms, frameworks, or libraries</u> for LLM evaluation (e.g., fast chat, vllm, huggingface) <u>that are used</u> .
33	<u>Repeat the experiment multiple times</u> to reduce the impact of <u>randomness</u> on the evaluation
34	Consider various <u>randomization strategies</u> (e.g., trying various temperature parameters) <u>to reduce the impact of parameter configuration on the evaluation</u> .
35	<u>Record the experimental process in detail</u> (e.g., parameter settings, running time, input/output pairs, etc.).

Figure 5: Guideline for Benchmark Evaluation

🔍 **Explanation** – Guidelines for benchmark evaluation focus on the rigorousness and reliability of the evaluation. HOW2BENCH provides 12 criteria for benchmark evaluation, as shown in Figure 5. It mainly focuses on the comprehensive evaluation processes for benchmarks involving LLMs. For evaluation design, it stresses the importance of assessing a sufficient and *representative range of LLMs* to ensure the benchmark’s applicability across various model families and configurations, both open and closed-source. Figure 29 and Figure 30 shows the distribution of numbers of LLMs studied and the most exercised LLMs.

Also, *prompting* has a direct impact on the qual-

ity of the LLMs’ output results (Wei et al., 2022; He et al., 2024a; Jin et al., 2024; Ye et al., 2023). As pointed out by a recent study, up to 40% performance gap could be observed in code translation when prompts vary (He et al., 2024b).

Additionally, *the experiment environment* is essential for reproducibility and transparency. Indeed, the hardware, software, and platform environments used during experiments might influence the outcomes (Ghosh, 2024). Furthermore, because of the nondeterministic nature of LLMs, experiments should be repeated, and randomization strategies should be used to mitigate the effects of randomness and parameter configuration biases. Lastly, *meticulously documented logs* of the experimental process are advised to facilitate transparency and reproducibility, detailing everything from parameter settings to the specific LLM pipelines such as vLLM (Kwon et al., 2023) used.

■ **Key Statistics** – Among the 274 benchmarks, 183 of them are evaluated over LLMs. According to our statistics (Figure 29), over 34% of the benchmarks were evaluated on fewer than 3 LLMs, with **11.48% benchmarks only evaluated on one LLM**. Such evaluation results can hardly be generalized to other LLMs. Furthermore, *more than half of the benchmarks studied fewer than 6 LLMs* ($51\% = (21 + 22 + 20 + 4 + 12 + 15)/183$).

☛ For reference, we listed the top 10 most studied LLM families in Figure 30. Among them, the GPT and CodeLlama series are the most extensively studied, accounting for 63% (116/183) and 33% (60/183), respectively. Under the constraints of time and available resources, it is beneficial to evaluate more representative LLMs.

The prompt quality also greatly impacts the LLM evaluation (He et al., 2024b). According to a recent study, up to 40% performance vary could be observed in code translation task (He et al., 2024b). So, carefully designing a prompt needs consideration. However, **73.3%** representative benchmarks (Appendix C) do not validate whether the prompts they used are well-designed (Appendix A.4). Similarly, though 94.9% benchmarks were evaluated in a zero-shot manner, only 21.2% benchmarks were evaluated under few-shot, 8.8% under Chain-of-Thought and 2.6% under RAG (Appendix A.4). However, as shown in Figure 34, 73.3% representative benchmarks (Appendix C) do not validate whether the prompt they used is well-designed.

Regarding the evaluation process, our statistics exposed that **only 35.4% of benchmark evalua-**

tions have been repeated (Appendix A.4). Also, regarding the transparency and matriculated documents, the observation is not optimistic – **Only 3.6% benchmarks provided their experiment environment. More than 50% of benchmarks did not provide reproducible instructions** such as prompts, examples for few-shot learning, or content for retrieval (Figure 39). **Less than half (42.7%) provide hyperparameters** such as temperature for reproduction.

▲ **Severity** – Over 60% of evaluations have not been repeated to eliminate the impact of randomness. **Only a few (less than 3.6%) provide the complete and necessary information required for reproducibility** such as prompts and environment.

4.4 Guideline for Evaluation Analysis

Phase 3. Benchmark Analysis	
36	Observe the difficulty of the benchmark, checking if the benchmark is too hard or too easy for LLMs (i.e., most LLMs score too high/low).
37	Consider whether the benchmark can distinguish the pros and cons of different LLMs.
38	If the experiment is repeated several times , consider the stability of the benchmark (i.e., whether the experimental results vary too much in the repeated experiments).
39	Analyze the correlation between the data and their score. For example, if there is a correlation between the data (such as similar difficulty and knowledge required), then the scores should also be correlated.
40	Compare the performance of LLMs on this benchmark with their performance on other related benchmarks .
41	Consider presenting the experiment results in an appropriate way (e.g., table, line graph, pie chart, etc.).
42	Consider presenting the experiment results clearly (e.g., distinguishable colors/labels/shapes, etc.).
43	Explain the experiment results.
44	Observe correlations via multiple perspectives from the experimental results (e.g., performance is correlated with model size or amount of context).
45	The analysis of the evaluation results will be inspiring (e.g., shed light on future direction, make actionable advice, etc.).

Figure 6: Guideline for Evaluation Analysis

🔍 **Explanation** – The analysis of the experiment results is expected to be objective and comprehensive, hopefully providing insights or actionable advice. So, we listed 10 criteria for the evaluation analysis phase, as shown in Figure 6. Regarding **the perspectives of analysis**, inspired by classic measurement theory (Suppes et al., 1962), we suggest four essential perspectives, including **difficulty** (whether a benchmark is appropriately challenging for LLMs), **stability** (whether the results are consistent through repeated trials), **differentiability** (whether benchmarks can differentiate the strengths and weaknesses of various LLMs), and **inspiration** (e.g., the correlations between the upper-/downstream coding tasks and LLM scores).

Moreover, effective **presentation of results** using clear visual and textual descriptions is recom-

mended to ensure the findings are understandable and actionable. The phase concludes with the suggestion to interpret and explain the results comprehensively, providing a basis for future research and application enhancements.

■ **Key Statistics** – Because experimental analysis is relatively subjective and cannot be obtained through mechanical scanning, we focus on 30 representative focus benchmarks (Appendix C), covering the highest cited and latest benchmarks in top five tasks. Figure 37 shows an example from CruxEval (Gu et al., 2024) where the experimental scores can hardly be read from the figures.

Also, **explaining experiment results** is crucial for other practitioners to understand what the outcomes mean in the context of the research questions. According to our statistics (Appendix A.5), 70% benchmarks have detailed explanations and analyses of their evaluation results, while still **30% have not**. Indeed, an explanation contributes to the body of knowledge by making it possible to understand and compare results with previous studies, promoting transparency within the community.

▲ **Severity** – The analysis of experimental data and the clarity of data presentation may receive less attention and worth consideration. Even in papers cited 1k+ times like MBPP (Austin et al., 2021), there are instances of **unclear evaluation analysis and display**.

4.5 Guideline for Benchmark Release

Phase 4. Benchmark Release	
46	Set the appropriate license for the benchmark.
47	Review the released benchmark or other artifacts to ensure they do NOT contain sensitive information (e.g., API keys, usernames, passwords, etc.).
48	review the released benchmark or other artifacts to ensure they do NOT contain toxicity information (e.g., abusive comments/identifiers).
49	Make sure the benchmark is open-accessible .
50	Make sure the test cases or reference data are open and accessible.
51	Provide prompts used in the experiment to ensure the experiments are reproducible.
52	Disclose the experimental environment (e.g., hardware, operating system, software version, framework platform) to ensure the reproducibility of the experiment .
53	Make the detailed experimental results public for verification.
54	Ensure the quality of the user manual such as README (e.g., it contains necessary benchmark introduction, executable scripts, etc.).
55	Provide convenient evaluation interfaces for the released benchmark (e.g., providing a command line interface, docker, etc.).

Figure 7: Guideline for Benchmark Release

🔍 **Explanation** – Finally, releasing a benchmark for open access also needs careful consideration. We offered 10 suggestions for this step, as shown in Figure 7, to highlight essential steps for public release preparation, emphasizing accessibility and ethical compliance. This includes setting an appro-

appropriate **license** to clarify usage rights, conducting a thorough review to **eliminate sensitive or harmful content** such as the API keys to access LLMs, the personal emails or toxic code comments (Miller et al., 2022) unless they are a part of the benchmark, and ensuring **transparency and reproducibility** by making all related materials openly available. **Detailed prompts** and clear descriptions of the experimental setup are advised to facilitate replication. Additionally, providing user manuals and evaluation interfaces is crucial for effective user engagement with the benchmark, enhancing its reliability and value for the research community.

■ **Key Statistics** – The final step involves the release of the benchmark. The fundamental requirement for releasing a benchmark is that it must be open-sourced. However, surprisingly, we observed that 5.1% of the benchmarks are only partially open-sourced (e.g., missing some subjects or tests), and **5.8% are not open-sourced at all** (e.g., links/web pages are no longer active). **19.3% have not properly set up the license**. Furthermore, prompts, which are necessary for reproducibility, are not disclosed in 52.6% of the benchmarks (Figure 39). Not to mention the lack of public information on experimental settings (Figure 32 and Figure 31) and experimental parameters (Figure 43). What is worse, 19.3% benchmarks do not setup licenses (Figure 44). The absence of licensing may lead to severe legal and ethical issues, potentially resulting in unauthorized use and distribution of proprietary technologies. Additionally, only 16.7% of the benchmarks make their logged experimental results publicly available (Appendix A.6).

▲ **Severity** – The release of existing benchmarks exhibits several issues. For example, over 10% of the benchmarks are either not open to public access or are only partially open-sourced. Only 47.4% of benchmarks are released with replicable prompts.

5 Human Study

To delve deeper into the integration of knowledge and action, we **surveyed 49 global researchers** in AI (42.6%) and SE (57.14%), as shown in Figure 50. Each participant had published at least one research paper, and about **half had constructed code-related benchmarks**. See Appendix B.

First, **all participants agreed** that having a checklist for benchmark construction would con-

tribute to the quality of the benchmark. 47/55 criteria in HOW2BENCH are deemed important by more 80% participants. Additionally, among the 21 participants who have constructed code-related benchmarks, **53 out of 55 criteria were deemed important by all benchmark developers**; only two criteria (criteria 3 and 4 in Section 4) were considered unimportant by a few individuals (3 and 2 participants, respectively). Additionally, we received two valuable suggestions that draw importance to recording **the time/monetary costs** of constructing the benchmark and conducting the experiments.

However, we also identified some **notable gaps in awareness**. First, regarding the **data preparation**, more than 15% of participants were not aware that the selection of data should consider the target scope of the evaluation set (i.e., the data must be representative), and 16% of participants were **unaware of the need for data denoising**. This oversight can significantly affect the validity and generalizability of experimental results, underscoring the importance of a comprehensive understanding of data handling for reliable research outcomes. Second, regarding **evaluation replicability and reliability**. Over 40% of participants believe that recording and publicizing the hardware and software environments, software versions, and libraries used in experiments is not important, with more than 20% still considering it unimportant despite already done so. This reveals **a significant lack of awareness** about the impact that experimental environments can have on **the reliability, reproducibility, and stability of evaluation results**. In fact, various studies have demonstrated that different experimental environments, parameters, and prompts can lead to substantial variations in outcomes (Xiao et al., 2024; Wang et al., 2019, 2023a).

6 Conclusion

This paper proposes a rigorous guideline consisting of 55 checklists covering the benchmark development lifecycle. After investigating over 270 code-related benchmarks, we exposed their merits and limitations and provided suggestions for improving them. Finally, our human study reveals the neglect of details that may affect the benchmark’s reliability. In the long run, HOW2BENCH helps to improve the overall quality of benchmarks in the community due to the propagation among benchmarks.

Limitations

This paper has two primary limitations that offer avenues for future research. First, *the collection of code-related benchmarks may be incomplete*. To minimize this limitation, we covered papers published over the last decade, and conducted multiple rounds of snowballing. Ultimately, we collected 274 benchmarks, which is comparable to the number included in recent surveys (Hou et al., 2023; Schäfer et al., 2024) in the field. Second, *the study involved substantial manual analysis*, which could lead to oversight and discrepancies in the statistical results. To mitigate this issue, we ensured that each benchmark was double-checked by at least two authors and underwent multiple rounds of iteration. Third, *the guidelines may not cover all the details*. Constructing a code-related benchmark involves numerous details, and some criteria are task-specific. To overcome this limitation, we iteratively refined the guidelines, interviewed practitioners, and tried to cover the entire benchmark development process as thoroughly as possible. Last, *the human study participants may exhibit subjectivity*. To address this limitation, we endeavored to include a broad range of practitioners and seasoned researchers with experience in both AI and SE, aiming for the least biased results possible.

References

Yash Agarwal, Devansh Batra, and Ganesh Bagler. 2020. [Building hierarchically disentangled language models for text generation with named entities](#). In *Proceedings of the 28th International Conference on Computational Linguistics, COLING 2020, Barcelona, Spain (Online), December 8-13, 2020*, pages 26–38. International Committee on Computational Linguistics.

Rajas Agashe, Srinivasan Iyer, and Luke Zettlemoyer. 2019. [Juice: A large scale distantly supervised dataset for open domain context-based code generation](#). In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing, EMNLP-IJCNLP 2019, Hong Kong, China, November 3-7, 2019*, pages 5435–5445. Association for Computational Linguistics.

Lakshya A. Agrawal, Aditya Kanade, Navin Goyal, Shuvendu K. Lahiri, and Sriram K. Rajamani. 2023. [Monitor-guided decoding of code lms with static analysis of repository context](#). In *Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023*.

Wasi Uddin Ahmad, Md Golam Rahman Tushar, Saikat Chakraborty, and Kai-Wei Chang. 2023. [AVATAR: A parallel corpus for java-python program translation](#). In *Findings of the Association for Computational Linguistics: ACL 2023, Toronto, Canada, July 9-14, 2023*, pages 2268–2281. Association for Computational Linguistics.

Miltiadis Allamanis, Henry Jackson-Flux, and Marc Brockschmidt. 2021. [Self-supervised bug detection and repair](#). In *Advances in Neural Information Processing Systems 34: Annual Conference on Neural Information Processing Systems 2021, NeurIPS 2021, December 6-14, 2021, virtual*, pages 27865–27876.

Uri Alon, Shaked Brody, Omer Levy, and Eran Yahav. 2019. [code2seq: Generating sequences from structured representations of code](#). In *7th International Conference on Learning Representations, ICLR 2019, New Orleans, LA, USA, May 6-9, 2019*. OpenReview.net.

Aida Amini, Saadia Gabriel, Shanchuan Lin, Rik Koncel-Kedziorski, Yejin Choi, and Hannaneh Hajishirzi. 2019. [MathQA: Towards interpretable math word problem solving with operation-based formalisms](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 2357–2367, Minneapolis, Minnesota. Association for Computational Linguistics.

Ben Athiwaratkun, Sanjay Krishna Gouda, Zijian Wang, Xiaopeng Li, Yuchen Tian, Ming Tan, Wasi Uddin Ahmad, Shiqi Wang, Qing Sun, Mingyue Shang, Sujan Kumar Gonugondla, Hantian Ding, Varun Kumar, Nathan Fulton, Arash Farahani, Siddhartha Jain, Robert Giaquinto, Haifeng Qian, Murali Krishna Ramanathan, Ramesh Nallapati, Baishakhi Ray, Parminder Bhatia, Sudipta Sengupta, Dan Roth, and Bing Xiang. 2022. [Multi-lingual evaluation of code generation models](#).

Jacob Austin, Augustus Odena, Maxwell Nye, Maarten Bosma, Henryk Michalewski, David Dohan, Ellen Jiang, Carrie Cai, Michael Terry, Quoc Le, et al. 2021. Program synthesis with large language models. *arXiv preprint arXiv:2108.07732*.

Hannah McLean Babe, Sydney Nguyen, Yangtian Zi, Arjun Guha, Molly Q. Feldman, and Carolyn Jane Anderson. 2024. [Studenteval: A benchmark of student-written prompts for large language models of code](#). In *Findings of the Association for Computational Linguistics, ACL 2024, Bangkok, Thailand and virtual meeting, August 11-16, 2024*, pages 8452–8474. Association for Computational Linguistics.

Ramakrishna Bairi, Atharv Sonwane, Aditya Kanade, Vageesh D. C., Arun Iyer, Suresh Parthasarathy, Sriram K. Rajamani, Balasubramanyan Ashok, and Shashank Shet. 2024. [Codeplan: Repository-level coding using llms and planning](#). *Proc. ACM Softw. Eng.*, 1(FSE):675–698.

- Antonio Valerio Miceli Barone and Rico Sennrich. 2017. [A parallel corpus of python functions and documentation strings for automated code documentation and code generation](#). In *Proceedings of the Eighth International Joint Conference on Natural Language Processing, IJCNLP 2017, Taipei, Taiwan, November 27 - December 1, 2017, Volume 2: Short Papers*, pages 314–319. Asian Federation of Natural Language Processing.
- Earl T Barr, Mark Harman, Phil McMinn, Muzammil Shahbaz, and Shin Yoo. 2014. The oracle problem in software testing: A survey. *IEEE transactions on software engineering*, 41(5):507–525.
- Berkay Berabi, Jingxuan He, Veselin Raychev, and Martin T. Vechev. 2021. [Tfix: Learning to fix coding errors with a text-to-text transformer](#). In *Proceedings of the 38th International Conference on Machine Learning, ICML 2021, 18-24 July 2021, Virtual Event*, volume 139 of *Proceedings of Machine Learning Research*, pages 780–791. PMLR.
- Egor Bogomolov, Aleksandra Eliseeva, Timur Galimzyanov, Evgeniy Glukhov, Anton Shapkin, Maria Tigina, Yaroslav Golubev, Alexander Kovrigin, Arie van Deursen, Maliheh Izadi, and Timofey Bryksin. 2024. [Long code arena: a set of benchmarks for long-context code models](#). *CoRR*, abs/2406.11612.
- Jialun Cao, Zhiyong Chen, Jiarong Wu, Shing-Chi Cheung, and Chang Xu. 2024a. [Can AI beat undergraduates in entry-level java assignments? benchmarking large language models on javabench](#). *CoRR*, abs/2406.12902.
- Jialun Cao, Wuqi Zhang, and Shing-Chi Cheung. 2024b. Concerned with data contamination? assessing countermeasures in code language model. *arXiv preprint arXiv:2403.16898*.
- Ruisheng Cao, Fangyu Lei, Haoyuan Wu, Jixuan Chen, Yeqiao Fu, Hongcheng Gao, Xinzhuang Xiong, Hanchong Zhang, Yuchen Mao, Wenjing Hu, Tianbao Xie, Hongshen Xu, Danyang Zhang, Sida Wang, Ruoxi Sun, Pengcheng Yin, Caiming Xiong, Ansong Ni, Qian Liu, Victor Zhong, Lu Chen, Kai Yu, and Tao Yu. 2024c. [Spider2-v: How far are multimodal agents from automating data science and engineering workflows?](#) *CoRR*, abs/2407.10956.
- Federico Cassano, John Gouwar, Daniel Nguyen, Sydney Nguyen, Luna Phipps-Costin, Donald Pinckney, Ming-Ho Yee, Yangtian Zi, Carolyn Jane Anderson, Molly Q Feldman, et al. 2022. Multipl-e: A scalable and extensible approach to benchmarking neural code generation. *arXiv preprint arXiv:2208.08227*.
- Saikat Chakraborty, Rahul Krishna, Yangruibo Ding, and Baishakhi Ray. 2022. [Deep learning based vulnerability detection: Are we there yet?](#) *IEEE Trans. Software Eng.*, 48(9):3280–3296.
- Shubham Chandel, Colin B Clement, Guillermo Serrato, and Neel Sundaresan. 2022. Training and evaluating a jupyter notebook data science assistant. *arXiv preprint arXiv:2201.12901*.
- Yupeng Chang, Xu Wang, Jindong Wang, Yuan Wu, Linyi Yang, Kaijie Zhu, Hao Chen, Xiaoyuan Yi, Cunxiang Wang, Yidong Wang, et al. 2024. A survey on evaluation of large language models. *ACM Transactions on Intelligent Systems and Technology*, 15(3):1–45.
- Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, Alex Ray, Raul Puri, Gretchen Krueger, Michael Petrov, Heidy Khlaaf, Girish Sastry, Pamela Mishkin, Brooke Chan, Scott Gray, Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz Kaiser, Mohammad Bavarian, Clemens Winter, Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebguss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Josh Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. 2021a. [Evaluating large language models trained on code](#).
- Xinyun Chen, Linyuan Gong, Alvin Cheung, and Dawn Song. 2021b. [Plotcoder: Hierarchical decoding for synthesizing visualization code in programmatic context](#). In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing, ACL/IJCNLP 2021, (Volume 1: Long Papers), Virtual Event, August 1-6, 2021*, pages 2169–2181. Association for Computational Linguistics.
- Yizheng Chen, Zhoujie Ding, Lamya Alowain, Xinyun Chen, and David A. Wagner. 2023. [Diversevul: A new vulnerable source code dataset for deep learning based vulnerability detection](#). In *Proceedings of the 26th International Symposium on Research in Attacks, Intrusions and Defenses, RAID 2023, Hong Kong, China, October 16-18, 2023*, pages 654–668. ACM.
- Jianbo Dai, Jianqiao Lu, Yunlong Feng, Rongju Ruan, Ming Cheng, Haochen Tan, and Zhijiang Guo. 2024. [MHPP: exploring the capabilities and limitations of language models beyond basic code generation](#). *CoRR*, abs/2405.11430.
- Xiang Deng, Ahmed Hassan Awadallah, Christopher Meek, Oleksandr Polozov, Huan Sun, and Matthew Richardson. 2021. [Structure-grounded pretraining for text-to-sql](#). In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, NAACL-HLT 2021, Online, June 6-11, 2021*, pages 1337–1350. Association for Computational Linguistics.
- Yangruibo Ding, Yanjun Fu, Omniyyah Ibrahim, Chawin Sitawarin, Xinyun Chen, Basel Alomair,

841	David A. Wagner, Baishakhi Ray, and Yizheng Chen.	pages 351–360. Association for Computational Lin-	898
842	2024a. Vulnerability detection with code language	guistics.	899
843	models: How far are we? <i>CoRR</i> , abs/2403.18624.		
844	Yangruibo Ding, Zijian Wang, Wasi Uddin Ahmad, Han-	Karl Fogel. 2005. <i>Producing open source software:</i>	900
845	tian Ding, Ming Tan, Nihal Jain, Murali Krishna Ra-	<i>How to run a successful free software project.</i> "	901
846	manathan, Ramesh Nallapati, Parminder Bhatia, Dan	O'Reilly Media, Inc.".	902
847	Roth, and Bing Xiang. 2023. Crosscodeeval: A di-		
848	verse and multilingual benchmark for cross-file code	Lingyue Fu, Huacan Chai, Shuang Luo, Kounianhua Du,	903
849	completion. In <i>Advances in Neural Information Pro-</i>	Weiming Zhang, Longteng Fan, Jiayi Lei, Renting	904
850	<i>cessing Systems 36: Annual Conference on Neural</i>	Rui, Jianghao Lin, Yuchen Fang, Yifan Liu, Jingkuan	905
851	<i>Information Processing Systems 2023, NeurIPS 2023,</i>	Wang, Siyuan Qi, Kangning Zhang, Weinan Zhang,	906
852	<i>New Orleans, LA, USA, December 10 - 16, 2023.</i>	and Yong Yu. 2023. Codeapex: A bilingual pro-	907
		gramming evaluation benchmark for large language	908
		models. <i>CoRR</i> , abs/2309.01940.	909
853	Yangruibo Ding, Zijian Wang, Wasi Uddin Ahmad,		
854	Murali Krishna Ramanathan, Ramesh Nallapati, Parm-	YanJun Fu, Ethan Baker, and Yizheng Chen. 2024. Con-	910
855	minder Bhatia, Dan Roth, and Bing Xiang. 2024b.	strained decoding for secure code generation. <i>CoRR</i> ,	911
856	Cocomic: Code completion by jointly modeling in-	abs/2405.00218.	912
857	file and cross-file context. In <i>Proceedings of the</i>		
858	<i>2024 Joint International Conference on Computa-</i>	Yujian Gan, Xinyun Chen, Qiuping Huang, and	913
859	<i>tional Linguistics, Language Resources and Evalua-</i>	Matthew Purver. 2022. Measuring and improving	914
860	<i>tion, LREC/COLING 2024, 20-25 May, 2024, Torino,</i>	compositional generalization in text-to-sql via com-	915
861	<i>Italy</i> , pages 3433–3445. ELRA and ICCL.	ponent alignment. In <i>Findings of the Association</i>	916
		<i>for Computational Linguistics: NAACL 2022, Seattle,</i>	917
862	Shihan Dou, Haoxiang Jia, Shenxi Wu, Huiyuan Zheng,	<i>WA, United States, July 10-15, 2022</i> , pages 831–843.	918
863	Weikang Zhou, Muling Wu, Mingxu Chai, Jessica	Association for Computational Linguistics.	919
864	Fan, Caishuang Huang, Yunbo Tao, et al. 2024.		
865	What's wrong with your code generated by large	Yujian Gan, Xinyun Chen, Qiuping Huang, Matthew	920
866	language models? an extensive study. <i>arXiv preprint</i>	Purver, John R. Woodward, Jinxia Xie, and Peng-	921
867	<i>arXiv:2407.06153.</i>	sheng Huang. 2021a. Towards robustness of text-to-	922
		sql models against synonym substitution. In <i>Proceed-</i>	923
868	Mingzhe Du, Anh Tuan Luu, Bin Ji, Qian Liu, and	<i>ings of the 59th Annual Meeting of the Association for</i>	924
869	See-Kiong Ng. 2024. Mercury: A code efficiency	<i>Computational Linguistics and the 11th International</i>	925
870	benchmark for code large language models. In <i>The</i>	<i>Joint Conference on Natural Language Processing,</i>	926
871	<i>Thirty-eight Conference on Neural Information Pro-</i>	<i>ACL/IJCNLP 2021, (Volume 1: Long Papers), Virtual</i>	927
872	<i>cessing Systems Datasets and Benchmarks Track.</i>	<i>Event, August 1-6, 2021</i> , pages 2505–2515. Associa-	928
		tion for Computational Linguistics.	929
873	Xueying Du, Mingwei Liu, Kaixin Wang, Hanlin Wang,		
874	Junwei Liu, Yixuan Chen, Jiayi Feng, Chaofeng	Yujian Gan, Xinyun Chen, and Matthew Purver. 2021b.	930
875	Sha, Xin Peng, and Yiling Lou. 2023a. Classe-	Exploring underexplored limitations of cross-domain	931
876	val: A manually-crafted benchmark for evaluat-	text-to-sql generalization. In <i>Proceedings of the 2021</i>	932
877	ing llms on class-level code generation. <i>Preprint,</i>	<i>Conference on Empirical Methods in Natural Lan-</i>	933
878	<i>arXiv:2308.01861.</i>	<i>guage Processing, EMNLP 2021, Virtual Event /</i>	934
		<i>Punta Cana, Dominican Republic, 7-11 November,</i>	935
879	Xueying Du, Mingwei Liu, Kaixin Wang, Hanlin Wang,	2021, pages 8926–8931. Association for Computa-	936
880	Junwei Liu, Yixuan Chen, Jiayi Feng, Chaofeng Sha,	tional Linguistics.	937
881	Xin Peng, and Yiling Lou. 2023b. Classeval: A		
882	manually-crafted benchmark for evaluating llms on	Luyu Gao, Aman Madaan, Shuyan Zhou, Uri Alon,	938
883	class-level code generation. <i>CoRR</i> , abs/2308.01861.	Pengfei Liu, Yiming Yang, Jamie Callan, and Gra-	939
		ham Neubig. 2023a. PAL: program-aided language	940
884	Aleksandra Eliseeva, Yaroslav Sokolov, Egor Bogo-	models. In <i>International Conference on Machine</i>	941
885	molv, Yaroslav Golubev, Danny Dig, and Timofey	<i>Learning, ICML 2023, 23-29 July 2023, Honolulu,</i>	942
886	Bryksin. 2023. From commit message generation	<i>Hawaii, USA</i> , volume 202 of <i>Proceedings of Machine</i>	943
887	to history-aware commit message completion. In	<i>Learning Research</i> , pages 10764–10799. PMLR.	944
888	<i>38th IEEE/ACM International Conference on Auto-</i>		
889	<i>mated Software Engineering, ASE 2023, Luxembourg,</i>	Zeyu Gao, Hao Wang, Yuchen Zhou, Wenyu Zhu, and	945
890	<i>September 11-15, 2023</i> , pages 723–735. IEEE.	Chao Zhang. 2023b. How far have we gone in vulner-	946
		ability detection using large language models. <i>CoRR</i> ,	947
891	Catherine Finegan-Dollak, Jonathan K. Kummerfeld,	abs/2311.12420.	948
892	Li Zhang, Karthik Ramanathan, Sesh Sadasivam, Rui		
893	Zhang, and Dragomir R. Radev. 2018. Improving	Spandan Garg, Roshanak Zilouchian Moghaddam,	949
894	text-to-sql evaluation methodology. In <i>Proceedings</i>	Colin B. Clement, Neel Sundaresan, and Chen	950
895	<i>of the 56th Annual Meeting of the Association for</i>	Wu. 2022. Deepperf: A deep learning-based ap-	951
896	<i>Computational Linguistics, ACL 2018, Melbourne,</i>	proach for improving software performance. <i>CoRR</i> ,	952
897	<i>Australia, July 15-20, 2018, Volume 1: Long Papers,</i>	abs/2206.13619.	953

954	Bijit Ghosh. 2024. Changing your gpu changes your llm behavior.	1008
955		1009
956	Shahriar Golchin and Mihai Surdeanu. 2023. Time travel in llms: Tracing data contamination in large language models. <i>CoRR</i> , abs/2308.08493.	1010
957		1011
958		1012
959	Jing Gong, Yanghui Wu, Linxi Liang, Zibin Zheng, and Yanlin Wang. 2024. Cosqa+: Enhancing code search dataset with matching code. <i>CoRR</i> , abs/2406.11589.	1013
960		1014
961		1015
962	Alex Gu, Baptiste Rozière, Hugh Leather, Armando Solar-Lezama, Gabriel Synnaeve, and Sida I Wang. 2024. Cruxeval: A benchmark for code reasoning, understanding and execution. <i>arXiv preprint arXiv:2401.03065</i> .	1016
963		1017
964		1018
965		1019
966		1020
967	Xiaodong Gu, Hongyu Zhang, and Sunghun Kim. 2018. Deep code search. In <i>Proceedings of the 40th International Conference on Software Engineering, ICSE 2018, Gothenburg, Sweden, May 27 - June 03, 2018</i> , pages 933–944. ACM.	1021
968		1022
969		1023
970		1024
971		1025
972	Jiawei Guo, Ziming Li, Xueling Liu, Kaijing Ma, Tianyu Zheng, Zhouliang Yu, Ding Pan, Yizhi Li, Ruibo Liu, Yue Wang, Shuyue Guo, Xingwei Qu, Xiang Yue, Ge Zhang, Wenhui Chen, and Jie Fu. 2024. Codeeditorbench: Evaluating code editing capability of large language models. <i>CoRR</i> , abs/2404.03543.	1026
973		1027
974		1028
975		1029
976		1030
977		1031
978	Rahul Gupta, Soham Pal, Aditya Kanade, and Shirish K. Shevade. 2017. Deepfix: Fixing common C language errors by deep learning. In <i>Proceedings of the Thirty-First AAAI Conference on Artificial Intelligence, February 4-9, 2017, San Francisco, California, USA</i> , pages 1345–1351. AAAI Press.	1032
979		1033
980		1034
981		1035
982		1036
983		1037
984	Nam Le Hai, Dung Manh Nguyen, and Nghi D. Q. Bui. 2024. On the impacts of contexts on repository-level code generation. <i>Preprint</i> , arXiv:2406.11927.	1038
985		1039
986		1040
987	Patrick Haller, Jonas Golde, and Alan Akbik. 2024. PECC: problem extraction and coding challenges. In <i>Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation, LREC/COLING 2024, 20-25 May, 2024, Torino, Italy</i> , pages 12690–12699. ELRA and ICCL.	1041
988		1042
989		1043
990		1044
991		1045
992		1046
993		1047
994	Yiyang Hao, Ge Li, Yongqiang Liu, Xiaowei Miao, He Zong, Siyuan Jiang, Yang Liu, and Wei He. 2022. Aixbench: A code generation benchmark dataset. <i>CoRR</i> , abs/2206.13179.	1048
995		1049
996		1050
997		1051
998	Md. Mahim Anjum Haque, Wasi Uddin Ahmad, Ismini Lourentzou, and Chris Brown. 2023. Fixeval: Execution-based evaluation of program fixes for programming problems. In <i>IEEE/ACM International Workshop on Automated Program Repair, APR@ICSE 2023, Melbourne, Australia, May 16, 2023</i> , pages 11–18. IEEE.	1052
999		1053
1000		1054
1001		1055
1002		1056
1003		1057
1004		1058
1005	Masum Hasan, Tanveer Muttaqueen, Abdullah Al Ish-tiaq, Kazi Sajeed Mehrab, Md. Mahim Anjum Haque, Tahmid Hasan, Wasi Uddin Ahmad, Anindya Iqbal, and Rifat Shahriyar. 2021. Codesc: A large code-description parallel dataset. In <i>Findings of the Association for Computational Linguistics: ACL/IJCNLP 2021, Online Event, August 1-6, 2021</i> , volume ACL/IJCNLP 2021 of <i>Findings of ACL</i> , pages 210–218. Association for Computational Linguistics.	1059
1006		1060
1007		1061
		1062
		1063
		1064
		1065
		1066
		1067
		1068
		1069
		1070
		1071
		1072
		1073
		1074
		1075
		1076
		1077
		1078
		1079
		1080
		1081
		1082
		1083
		1084
		1085
		1086
		1087
		1088
		1089
		1090
		1091
		1092
		1093
		1094
		1095
		1096
		1097
		1098
		1099
		1100
		1101
		1102
		1103
		1104
		1105
		1106
		1107
		1108
		1109
		1110
		1111
		1112
		1113
		1114
		1115
		1116
		1117
		1118
		1119
		1120
		1121
		1122
		1123
		1124
		1125
		1126
		1127
		1128
		1129
		1130
		1131
		1132
		1133
		1134
		1135
		1136
		1137
		1138
		1139
		1140
		1141
		1142
		1143
		1144
		1145
		1146
		1147
		1148
		1149
		1150
		1151
		1152
		1153
		1154
		1155
		1156
		1157
		1158
		1159
		1160
		1161
		1162
		1163
		1164
		1165
		1166
		1167
		1168
		1169
		1170
		1171
		1172
		1173
		1174
		1175
		1176
		1177
		1178
		1179
		1180
		1181
		1182
		1183
		1184
		1185
		1186
		1187
		1188
		1189
		1190
		1191
		1192
		1193
		1194
		1195
		1196
		1197
		1198
		1199
		1200

1065	Yang Hu, Umair Z. Ahmed, Sergey Mechtaev, Ben	Nan Jiang, Kevin Liu, Thibaud Lutellier, and Lin Tan.	1122
1066	Leong, and Abhik Roychoudhury. 2019. Re-	2023. Impact of code language models on automated	1123
1067	factoring based program repair applied to program-	program repair . In <i>45th IEEE/ACM International</i>	1124
1068	ming assignments . In <i>2019 34th IEEE/ACM Interna-</i>	<i>Conference on Software Engineering, ICSE 2023,</i>	1125
1069	<i>tional Conference on Automated Software Engineer-</i>	<i>Melbourne, Australia, May 14-20, 2023, pages 1430–</i>	1126
1070	<i>ing (ASE), pages 388–398.</i>	1442. IEEE.	1127
1071	Dong HUANG, Yuhao QING, Weiyi Shang, Heming	Mingsheng Jiao, Tingrui Yu, Xuan Li, Guanjie Qiu,	1128
1072	Cui, and Jie Zhang. 2024. Effibench: Benchmarking	Xiaodong Gu, and Beijun Shen. 2023. On the eval-	1129
1073	the efficiency of automatically generated code . In	uation of neural code translation: Taxonomy and	1130
1074	<i>The Thirty-eight Conference on Neural Information</i>	benchmark . In <i>38th IEEE/ACM International Con-</i>	1131
1075	<i>Processing Systems Datasets and Benchmarks Track.</i>	<i>ference on Automated Software Engineering, ASE</i>	1132
1076	Junjie Huang, Chenglong Wang, Jipeng Zhang, Cong	2023, <i>Luxembourg, September 11-15, 2023, pages</i>	1133
1077	Yan, Haotian Cui, Jeevana Priya Inala, Colin B.	1529–1541. IEEE.	1134
1078	Clement, Nan Duan, and Jianfeng Gao. 2022.	Carlos E Jimenez, John Yang, Alexander Wettig,	1135
1079	Execution-based evaluation for data science code	Shunyu Yao, Kexin Pei, Ofir Press, and Karthik R	1136
1080	generation models . <i>CoRR</i> , abs/2211.09374.	Narasimhan. 2024. SWE-bench: Can language mod-	1137
1081	Yiming Huang, Zhenghao Lin, Xiao Liu, Yeyun Gong,	els resolve real-world github issues? In <i>The Twelfth</i>	1138
1082	Shuai Lu, Fangyu Lei, Yaobo Liang, Yelong Shen,	<i>International Conference on Learning Representa-</i>	1139
1083	Chen Lin, Nan Duan, and Weizhu Chen. 2024.	<i>tions</i> .	1140
1084	Competition-level problems are effective LLM eval-	Matthew Jin, Syed Shahriar, Michele Tufano, Xin	1141
1085	uators . In <i>Findings of the Association for Computa-</i>	Shi, Shuai Lu, Neel Sundaresan, and Alexey Svy-	1142
1086	<i>tional Linguistics, ACL 2024, Bangkok, Thailand and</i>	atkovskiy. 2023. Inferfix: End-to-end program repair	1143
1087	<i>virtual meeting, August 11-16, 2024, pages 13526–</i>	with llms . In <i>Proceedings of the 31st ACM Joint Eu-</i>	1144
1088	13544. Association for Computational Linguistics.	<i>ropean Software Engineering Conference and Sym-</i>	1145
1089	Hamel Husain, Ho-Hsiang Wu, Tiferet Gazit, Miltiadis	<i>posium on the Foundations of Software Engineering,</i>	1146
1090	Allamanis, and Marc Brockschmidt. 2019. Code-	<i>ESEC/FSE 2023, San Francisco, CA, USA, Decem-</i>	1147
1091	searchnet challenge: Evaluating the state of semantic	<i>ber 3-9, 2023, pages 1646–1656. ACM.</i>	1148
1092	code search . <i>CoRR</i> , abs/1909.09436.	Mingyu Jin, Qinkai Yu, Dong Shu, Haiyan Zhao,	1149
1093	Shima Imani, Liang Du, and Harsh Shrivastava. 2023.	Wenyue Hua, Yanda Meng, Yongfeng Zhang, and	1150
1094	Mathprompter: Mathematical reasoning using large	Mengnan Du. 2024. The impact of reasoning step	1151
1095	language models . <i>Preprint</i> , arXiv:2303.05398.	length on large language models. <i>arXiv preprint</i>	1152
1096	Marko Ivanković, Goran Petrović, René Just, and Gor-	<i>arXiv:2401.04925</i> .	1153
1097	don Fraser. 2019. Code coverage at google . In	René Just, Darioush Jalali, and Michael D. Ernst. 2014.	1154
1098	<i>Proceedings of the 2019 27th ACM Joint Meeting</i>	Defects4j: a database of existing faults to enable con-	1155
1099	<i>on European Software Engineering Conference and</i>	trolled testing studies for java programs . In <i>Interna-</i>	1156
1100	<i>Symposium on the Foundations of Software Engineer-</i>	<i>tional Symposium on Software Testing and Analysis,</i>	1157
1101	<i>ing, ESEC/FSE 2019, page 955–963, New York, NY,</i>	<i>ISSTA '14, San Jose, CA, USA - July 21 - 26, 2014,</i>	1158
1102	USA. Association for Computing Machinery.	pages 437–440. ACM.	1159
1103	Srinivasan Iyer, Ioannis Konstas, Alvin Cheung, and	Mohammad Abdullah Matin Khan, M. Saiful Bari,	1160
1104	Luke Zettlemoyer. 2016. Summarizing source code	Xuan Do Long, Weishi Wang, Md. Rizwan Parvez,	1161
1105	using a neural attention model . In <i>Proceedings of the</i>	and Shafiq Joty. 2024. Xcodeeval: An execution-	1162
1106	<i>54th Annual Meeting of the Association for Computa-</i>	based large scale multilingual multitask benchmark	1163
1107	<i>tional Linguistics, ACL 2016, August 7-12, 2016,</i>	for code understanding, generation, translation and	1164
1108	<i>Berlin, Germany, Volume 1: Long Papers. The Asso-</i>	retrieval . In <i>Proceedings of the 62nd Annual Meeting</i>	1165
1109	ciation for Computer Linguistics.	<i>of the Association for Computational Linguistics (Vol-</i>	1166
1110	Srinivasan Iyer, Ioannis Konstas, Alvin Cheung, and	<i>ume 1: Long Papers), ACL 2024, Bangkok, Thailand,</i>	1167
1111	Luke Zettlemoyer. 2018. Mapping language to code	<i>August 11-16, 2024, pages 6766–6805. Association</i>	1168
1112	in programmatic context . In <i>Proceedings of the 2018</i>	<i>for Computational Linguistics.</i>	1169
1113	<i>Conference on Empirical Methods in Natural Lan-</i>	Rahul Kumar, Amar Raja Dibbu, Shrutendra Harsola,	1170
1114	<i>guage Processing, pages 1643–1652, Brussels, Bel-</i>	Vignesh Subrahmaniam, and Ashutosh Modi. 2024.	1171
1115	gium. Association for Computational Linguistics.	Booksql: A large scale text-to-sql dataset for account-	1172
1116	Naman Jain, King Han, Alex Gu, Wen-Ding Li, Fanjia	ing domain . In <i>Proceedings of the 2024 Conference</i>	1173
1117	Yan, Tianjun Zhang, Sida Wang, Armando Solar-	<i>of the North American Chapter of the Association</i>	1174
1118	Lezama, Koushik Sen, and Ion Stoica. 2024. Live-	<i>for Computational Linguistics: Human Language</i>	1175
1119	codebench: Holistic and contamination free eval-	<i>Technologies (Volume 1: Long Papers), NAACL 2024,</i>	1176
1120	uation of large language models for code . <i>CoRR</i> ,	<i>Mexico City, Mexico, June 16-21, 2024, pages 497–</i>	1177
1121	abs/2403.07974.	516. Association for Computational Linguistics.	1178

1179	Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying	Fang, Lanshen Wang, Jiazheng Ding, Xuanming	1237
1180	Sheng, Lianmin Zheng, Cody Hao Yu, Joseph E.	Zhang, Yuqi Zhu, Yihong Dong, Zhi Jin, Binhua	1238
1181	Gonzalez, Hao Zhang, and Ion Stoica. 2023. Effi-	Li, Fei Huang, and Yongbin Li. 2024a. DevEval: A	1239
1182	cient memory management for large language model	Manually-Annotated Code Generation Benchmark	1240
1183	serving with pagedattention. In <i>Proceedings of the</i>	Aligned with Real-World Code Repositories . <i>arXiv</i>	1241
1184	<i>ACM SIGOPS 29th Symposium on Operating Systems</i>	<i>preprint</i> . ArXiv:2405.19856 [cs].	1242
1185	<i>Principles</i> .		
1186	Beck LaBash, August Rosedale, Alex Reents, Lucas	Jinyang Li, Binyuan Hui, Ge Qu, Jiayi Yang, Binhua Li,	1243
1187	Negritto, and Colin Wiel. 2024. RES-Q: evaluating	Bowen Li, Bailin Wang, Bowen Qin, Ruiying Geng,	1244
1188	code-editing large language model systems at the	Nan Huo, Xuanhe Zhou, Chenhao Ma, Guoliang	1245
1189	repository scale . <i>CoRR</i> , abs/2406.16801.	Li, Kevin Chen-Chuan Chang, Fei Huang, Reynold	1246
1190	Yuhang Lai, Chengxi Li, Yiming Wang, Tianyi Zhang,	Cheng, and Yongbin Li. 2023a. Can LLM already	1247
1191	Ruiqi Zhong, Luke Zettlemoyer, Wen-Tau Yih,	serve as A database interface? A big bench for large-	1248
1192	Daniel Fried, Sida I. Wang, and Tao Yu. 2023. DS-	scale database grounded text-to-sqls . In <i>Advances in</i>	1249
1193	1000: A natural and reliable benchmark for data sci-	<i>Neural Information Processing Systems 36: Annual</i>	1250
1194	ence code generation . In <i>International Conference</i>	<i>Conference on Neural Information Processing Sys-</i>	1251
1195	<i>on Machine Learning, ICML 2023, 23-29 July 2023,</i>	<i>tems 2023, NeurIPS 2023, New Orleans, LA, USA,</i>	1252
1196	<i>Honolulu, Hawaii, USA</i> , volume 202 of <i>Proceedings</i>	<i>December 10 - 16, 2023</i> .	1253
1197	<i>of Machine Learning Research</i> , pages 18319–18345.		
1198	PMLR.	Kaixin Li, Qisheng Hu, James Xu Zhao, Hui Chen,	1254
1199	Claire Le Goues, Neal J. Holtschulte, Edward K. Smith,	Yuxi Xie, Tiedong Liu, Michael Shieh, and Junxian	1255
1200	Yuriy Brun, Premkumar T. Devanbu, Stephanie For-	He. 2024b. Instructcoder: Instruction tuning large	1256
1201	rest, and Westley Weimer. 2015. The manybugs and	language models for code editing . In <i>Proceedings</i>	1257
1202	introclass benchmarks for automated repair of C pro-	<i>of the 62nd Annual Meeting of the Association for</i>	1258
1203	grams . <i>IEEE Trans. Software Eng.</i> , 41(12):1236–	<i>Computational Linguistics, ACL 2024 - Student Re-</i>	1259
1204	1256.	<i>search Workshop, Bangkok, Thailand, August 11-16,</i>	1260
1205	Alexander LeClair, Siyuan Jiang, and Collin McMil-	2024, pages 50–70. Association for Computational	1261
1206	lan. 2019. A neural model for generating natural	Linguistics.	1262
1207	language summaries of program subroutines . In <i>Pro-</i>	Kaixin Li, Yuchen Tian, Qisheng Hu, Ziyang Luo, and	1263
1208	<i>ceedings of the 41st International Conference on</i>	Jing Ma. 2024c. Mmcode: Evaluating multi-modal	1264
1209	<i>Software Engineering, ICSE 2019, Montreal, QC,</i>	code large language models with visually rich pro-	1265
1210	<i>Canada, May 25-31, 2019</i> , pages 795–806. IEEE /	gramming problems . <i>CoRR</i> , abs/2404.09486.	1266
1211	ACM.		
1212	Changyoon Lee, Yeon Seonwoo, and Alice Oh. 2022.	Linyi Li, Shijie Geng, Zhenwen Li, Yibo He, Hao Yu,	1267
1213	CSIQA: A dataset for assisting code-based question	Ziyue Hua, Guanghan Ning, Siwei Wang, Tao Xie,	1268
1214	answering in an introductory programming course .	and Hongxia Yang. 2024d. Infibench: Evaluating	1269
1215	In <i>Proceedings of the 2022 Conference of the North</i>	the question-answering capabilities of code large lan-	1270
1216	<i>American Chapter of the Association for Computa-</i>	guage models . <i>Preprint</i> , arXiv:2404.07940.	1271
1217	<i>tional Linguistics: Human Language Technologies,</i>		
1218	<i>NAACL 2022, Seattle, WA, United States, July 10-15,</i>	Rongao Li, Jie Fu, Bo-Wen Zhang, Tao Huang, Zhihong	1272
1219	2022, pages 2026–2040. Association for Computa-	Sun, Chen Lyu, Guang Liu, Zhi Jin, and Ge Li. 2023b.	1273
1220	tional Linguistics.	TACO: topics in algorithmic code generation dataset .	1274
1221	Chia-Hsuan Lee, Oleksandr Polozov, and Matthew	<i>CoRR</i> , abs/2312.14852.	1275
1222	Richardson. 2021. Kaggledbqa: Realistic evalua-	Yujia Li, David Choi, Junyoung Chung, Nate Kush-	1276
1223	tion of text-to-sql parsers . In <i>Proceedings of the 59th</i>	man, Julian Schrittwieser, Rémi Leblond, Tom Ec-	1277
1224	<i>Annual Meeting of the Association for Computational</i>	cles, James Keeling, Felix Gimeno, Agustin Dal	1278
1225	<i>Linguistics and the 11th International Joint Confer-</i>	Lago, Thomas Hubert, Peter Choy, Cyprien de Mas-	1279
1226	<i>ence on Natural Language Processing, ACL/IJCNLP</i>	son d’Autume, Igor Babuschkin, Xinyun Chen, Po-	1280
1227	<i>2021, (Volume 1: Long Papers), Virtual Event, Au-</i>	Sen Huang, Johannes Welbl, Sven Gowal, Alexey	1281
1228	<i>gust 1-6, 2021</i> , pages 2261–2273. Association for	Cherepanov, James Molloy, Daniel J. Mankowitz,	1282
1229	Computational Linguistics.	Esme Sutherland Robson, Pushmeet Kohli, Nando	1283
1230	Gyubok Lee, Hyeonji Hwang, Seongsu Bae, Yeonsu	de Freitas, Koray Kavukcuoglu, and Oriol Vinyals.	1284
1231	Kwon, Woncheol Shin, Seongjun Yang, Minjoon Seo,	2022. Competition-level code generation with alpha-	1285
1232	Jong-Yeup Kim, and Edward Choi. 2023. EHRSQL:	code . <i>Science</i> , 378(6624):1092–1097.	1286
1233	A practical text-to-sql benchmark for electronic	Zehan Li, Jianfei Zhang, Chuantao Yin, Yuanxin	1287
1234	health records . <i>CoRR</i> , abs/2301.07695.	Ouyang, and Wenge Rong. 2024e. Procqa: A large-	1288
1235	Jia Li, Ge Li, Yunfei Zhao, Yongmin Li, Huanyu	scale community-based programming question	1289
1236	Liu, Hao Zhu, Lecheng Wang, Kaibo Liu, Zheng	answering dataset for code search . In <i>Proceedings of</i>	1290
		<i>the 2024 Joint International Conference on Computa-</i>	1291
		<i>tional Linguistics, Language Resources and Evalua-</i>	1292
		<i>tion, LREC/COLING 2024, 20-25 May, 2024, Torino,</i>	1293
		<i>Italy</i> , pages 13057–13067. ELRA and ICCL.	1294

1295	Zhen Li, Deqing Zou, Shouhuai Xu, Hai Jin, Yawei	Lingming Zhang. 2024a. Repoqa: Evaluating long	1352
1296	Zhu, Zhaoxuan Chen, Sujuan Wang, and Jialai	context code understanding . <i>CoRR</i> , abs/2406.06025.	1353
1297	Wang. 2018a. Sysevr: A framework for using deep		
1298	learning to detect software vulnerabilities . <i>CoRR</i> ,	Jiawei Liu, Chunqiu Steven Xia, Yuyao Wang, and Ling-	1354
1299	abs/1807.06756.	ming Zhang. 2023a. Is your code generated by chat-	1355
		GPT really correct? rigorous evaluation of large lan-	1356
1300	Zhen Li, Deqing Zou, Shouhuai Xu, Xinyu Ou, Hai	guage models for code generation . In <i>Thirty-seventh</i>	1357
1301	Jin, Sujuan Wang, Zhijun Deng, and Yuyi Zhong.	<i>Conference on Neural Information Processing Sys-</i>	1358
1302	2018b. Vuldeepecker: A deep learning-based system	<i>tems</i> .	1359
1303	for vulnerability detection . In <i>25th Annual Network</i>		
1304	<i>and Distributed System Security Symposium, NDSS</i>	Jiawei Liu, Chunqiu Steven Xia, Yuyao Wang, and Ling-	1360
1305	<i>2018, San Diego, California, USA, February 18-21,</i>	ming Zhang. 2023b. Is your code generated by chat-	1361
1306	<i>2018</i> . The Internet Society.	gpt really correct? rigorous evaluation of large lan-	1362
		guage models for code generation . In <i>Advances in</i>	1363
1307	Dianshu Liao, Shidong Pan, Xiaoyu Sun, Xiaoxue	<i>Neural Information Processing Systems 36: Annual</i>	1364
1308	Ren, Qing Huang, Zhenchang Xing, Huan Jin, and	<i>Conference on Neural Information Processing Sys-</i>	1365
1309	Qinying Li. 2024. A 3-codgen: A repository-level	<i>tems 2023, NeurIPS 2023, New Orleans, LA, USA,</i>	1366
1310	code generation framework for code reuse with local-	<i>December 10 - 16, 2023</i> .	1367
1311	aware, global-aware, and third-party-library-aware.		
1312	<i>IEEE Transactions on Software Engineering</i> .	Shangqing Liu, Yu Chen, Xiaofei Xie, Jing Kai Siow,	1368
		and Yang Liu. 2021. Retrieval-augmented generation	1369
1313	Derrick Lin, James Koppel, Angela Chen, and Armando	for code summarization via hybrid GNN . In <i>9th In-</i>	1370
1314	Solar-Lezama. 2017. Quixbugs: a multi-lingual	<i>ternational Conference on Learning Representations,</i>	1371
1315	program repair benchmark set based on the quixey	<i>ICLR 2021, Virtual Event, Austria, May 3-7, 2021.</i>	1372
1316	challenge . In <i>Proceedings Companion of the 2017</i>	OpenReview.net.	1373
1317	<i>ACM SIGPLAN International Conference on Systems,</i>		
1318	<i>Programming, Languages, and Applications: Soft-</i>	Shangqing Liu, Cuiyun Gao, Sen Chen, Lun Yiu Nie,	1374
1319	<i>ware for Humanity, SPLASH 2017, Vancouver, BC,</i>	and Yang Liu. 2022. ATOM: commit message gener-	1375
1320	<i>Canada, October 23 - 27, 2017</i> , pages 55–56. ACM.	ation based on abstract syntax tree and hybrid ranking .	1376
		<i>IEEE Trans. Software Eng.</i> , 48(5):1800–1817.	1377
1321	Guanjun Lin, Wei Xiao, Jun Zhang, and Yang Xiang.		
1322	2019. Deep learning-based vulnerable function de-	Tianyang Liu, Canwen Xu, and Julian J. McAuley.	1378
1323	tection: A benchmark . In <i>Information and Commu-</i>	2024b. Repobench: Benchmarking repository-level	1379
1324	<i>nications Security - 21st International Conference,</i>	code auto-completion systems . In <i>The Twelfth In-</i>	1380
1325	<i>ICICS 2019, Beijing, China, December 15-17, 2019,</i>	<i>ternational Conference on Learning Representations,</i>	1381
1326	<i>Revised Selected Papers</i> , volume 11999 of <i>Lecture</i>	<i>ICLR 2024, Vienna, Austria, May 7-11, 2024</i> . Open-	1382
1327	<i>Notes in Computer Science</i> , pages 219–232. Springer.	Review.net.	1383
1328	Guanjun Lin, Jun Zhang, Wei Luo, Lei Pan, Olivier Y.	Yu Liu, Lang Gao, Mingxin Yang, Yu Xie, Ping Chen,	1384
1329	de Vel, Paul Montague, and Yang Xiang. 2021.	Xiaojin Zhang, and Wei Chen. 2024c. Vuldetect-	1385
1330	Software vulnerability discovery via learning multi-	bench: Evaluating the deep capability of vulnera-	1386
1331	domain knowledge bases . <i>IEEE Trans. Dependable</i>	bility detection with large language models . <i>CoRR</i> ,	1387
1332	<i>Secur. Comput.</i> , 18(5):2469–2485.	abs/2406.07595.	1388
1333	Guanjun Lin, Jun Zhang, Wei Luo, Lei Pan, Yang Xiang,	Yuliang Liu, Xiangru Tang, Zefan Cai, Junjie Lu,	1389
1334	Olivier Y. de Vel, and Paul Montague. 2018. Cross-	Yichi Zhang, Yanjun Shao, Zexuan Deng, Helan Hu,	1390
1335	project transfer representation learning for vulnerable	Zengxian Yang, Kaikai An, Ruijun Huang, Shuzheng	1391
1336	function discovery . <i>IEEE Trans. Ind. Informatics</i> ,	Si, Sheng Chen, Haozhe Zhao, Zhengliang Li, Liang	1392
1337	14(7):3289–3297.	Chen, Yiming Zong, Yan Wang, Tianyu Liu, Zhi-	1393
		wei Jiang, Baobao Chang, Yujia Qin, Wangchunshu	1394
1338	Xiang Ling, Lingfei Wu, Saizhuo Wang, Gaoning Pan,	Zhou, Yilun Zhao, Arman Cohan, and Mark Ger-	1395
1339	Tengfei Ma, Fangli Xu, Alex X. Liu, Chunming Wu,	stein. 2023c. ML-bench: Large language models	1396
1340	and Shouling Ji. 2021. Deep graph matching and	leverage open-source libraries for machine learning	1397
1341	searching for semantic code retrieval . <i>ACM Trans.</i>	tasks . <i>CoRR</i> , abs/2311.09835.	1398
1342	<i>Knowl. Discov. Data</i> , 15(5):88:1–88:21.		
		Zhongxin Liu, Xin Xia, Ahmed E. Hassan, David Lo,	1399
1343	Chenxiao Liu and Xiaojun Wan. 2021. Codeqa: A	Zhenchang Xing, and Xinyu Wang. 2018. Neural-	1400
1344	question answering dataset for source code compre-	machine-translation-based commit message gener-	1401
1345	hension . In <i>Findings of the Association for Computa-</i>	ation: how far are we? In <i>Proceedings of the</i>	1402
1346	<i>tional Linguistics: EMNLP 2021, Virtual Event /</i>	<i>33rd ACM/IEEE International Conference on Auto-</i>	1403
1347	<i>Punta Cana, Dominican Republic, 16-20 November,</i>	<i>mated Software Engineering, ASE 2018, Montpellier,</i>	1404
1348	<i>2021</i> , pages 2618–2632. Association for Computa-	<i>France, September 3-7, 2018</i> , pages 373–384. ACM.	1405
1349	<i>tional Linguistics</i> .		
		Shuai Lu, Daya Guo, Shuo Ren, Junjie Huang, Alexey	1406
1350	Jiawei Liu, Jia Le Tian, Vijay Daita, Yuxiang Wei,	Svyatkovskiy, Ambrosio Blanco, Colin B. Clement,	1407
1351	Yifeng Ding, Yuhan Katherine Wang, Jun Yang, and		

1408	Dawn Drain, Daxin Jiang, Duyu Tang, Ge Li, Li-	Daniel Nichols, Joshua Hoke Davis, Zhaojun Xie, Ar-	1465
1409	dong Zhou, Linjun Shou, Long Zhou, Michele Tu-	jun Rajaram, and Abhinav Bhatele. 2024. Can large	1466
1410	fano, Ming Gong, Ming Zhou, Nan Duan, Neel Sun-	language models write parallel code? In <i>Proceed-</i>	1467
1411	daresan, Shao Kun Deng, Shengyu Fu, and Shujie	<i>ings of the 33rd International Symposium on High-</i>	1468
1412	Liu. 2021. Codexglue: A machine learning bench-	<i>Performance Parallel and Distributed Computing,</i>	1469
1413	mark dataset for code understanding and generation.	<i>HPDC 2024, Pisa, Italy, June 3-7, 2024</i> , pages 281–	1470
1414	In <i>Proceedings of the Neural Information Process-</i>	294. ACM.	1471
1415	<i>ing Systems Track on Datasets and Benchmarks 1,</i>		
1416	<i>NeurIPS Datasets and Benchmarks 2021, December</i>	Pengyu Nie, Rahul Banerjee, Junyi Jessy Li, Raymond J.	1472
1417	<i>2021, virtual.</i>	Mooney, and Milos Gligoric. 2023. Learning deep	1473
		semantics for test completion. In <i>45th IEEE/ACM</i>	1474
1418	Rabee Sohail Malik, Jibesh Patra, and Michael Pradel.	<i>International Conference on Software Engineering,</i>	1475
1419	2019. NI2type: inferring javascript function types	<i>ICSE 2023, Melbourne, Australia, May 14-20, 2023,</i>	1476
1420	from natural language information. In <i>Proceedings</i>	pages 2111–2123. IEEE.	1477
1421	<i>of the 41st International Conference on Software En-</i>		
1422	<i>gineering, ICSE 2019, Montreal, QC, Canada, May</i>	Erik Nijkamp, Bo Pang, Hiroaki Hayashi, Lifu Tu, Huan	1478
1423	<i>25-31, 2019</i> , pages 304–315. IEEE / ACM.	Wang, Yingbo Zhou, Silvio Savarese, and Caiming	1479
		Xiong. 2022. Codegen: An open large language	1480
1424	Courtney Miller, Sophie Cohen, Daniel Klug, Bogdan	model for code with multi-turn program synthesis.	1481
1425	Vasilescu, and Christian KaUstner. 2022. "did you	<i>arXiv preprint arXiv:2203.13474.</i>	1482
1426	miss my comment or what?": understanding toxicity		
1427	in open source discussions. In <i>Proceedings of the</i>	Georgios Nikitopoulos, Konstantina Dritsa, Panos	1483
1428	<i>44th International Conference on Software Engineer-</i>	Louridas, and Dimitris Mitropoulos. 2021. Crossvul:	1484
1429	<i>ing, ICSE '22</i> , page 710–722, New York, NY, USA.	a cross-language vulnerability dataset with commit	1485
1430	Association for Computing Machinery.	data. In <i>ESEC/FSE '21: 29th ACM Joint Euro-</i>	1486
		<i>pean Software Engineering Conference and Sympo-</i>	1487
1431	Amir M. Mir, Evaldas Latoskinas, Sebastian Proksch,	<i>sium on the Foundations of Software Engineering,</i>	1488
1432	and Georgios Gousios. 2022. Type4py: Practical	<i>Athens, Greece, August 23-28, 2021</i> , pages 1565–	1489
1433	deep similarity learning-based type inference for	1569. ACM.	1490
1434	python. In <i>44th IEEE/ACM 44th International Con-</i>		
1435	<i>ference on Software Engineering, ICSE 2022, Pitts-</i>	Wonseok Oh and Hakjoo Oh. 2022. Pyter: effective pro-	1491
1436	<i>burgh, PA, USA, May 25-27, 2022</i> , pages 2241–2252.	gram repair for python type errors. In <i>Proceedings of</i>	1492
1437	ACM.	<i>the 30th ACM Joint European Software Engineering</i>	1493
		<i>Conference and Symposium on the Foundations of</i>	1494
1438	Lili Mou, Ge Li, Lu Zhang, Tao Wang, and Zhi Jin. 2016.	<i>Software Engineering, ESEC/FSE 2022, Singapore,</i>	1495
1439	Convolutional neural networks over tree structures	<i>Singapore, November 14-18, 2022</i> , pages 922–934.	1496
1440	for programming language processing. In <i>Proceed-</i>	ACM.	1497
1441	<i>ings of the Thirtieth AAAI Conference on Artificial</i>		
1442	<i>Intelligence, February 12-17, 2016, Phoenix, Ari-</i>	Rangeet Pan, Ali Reza Ibrahimzada, Rahul Krishna,	1498
1443	<i>zona, USA</i> , pages 1287–1293. AAAI Press.	Divya Sankar, Lambert Pouguem Wassi, Michele	1499
		Merler, Boris Sobolev, Raju Pavuluri, Saurabh Sinha,	1500
1444	Hussein Mozannar, Valerie Chen, Mohammed Alsobay,	and Reyhaneh Jabbarvand. 2024. Lost in transla-	1501
1445	Subhro Das, Sebastian Zhao, Dennis Wei, Manish	tion: A study of bugs introduced by large language	1502
1446	Nagireddy, Prasanna Sattigeri, Ameet Talwalkar, and	models while translating code. In <i>Proceedings of the</i>	1503
1447	David A. Sontag. 2024. The realhumaneval: Eval-	<i>IEEE/ACM 46th International Conference on Soft-</i>	1504
1448	uating large language models' abilities to support	<i>ware Engineering, ICSE '24, New York, NY, USA.</i>	1505
1449	programmers. <i>CoRR</i> , abs/2404.02806.	Association for Computing Machinery.	1506
1450	Niklas Muennighoff, Qian Liu, Armel Randy Ze-	Shishir G. Patil, Tianjun Zhang, Xin Wang, and	1507
1451	baze, Qinkai Zheng, Binyuan Hui, Terry Yue Zhuo,	Joseph E. Gonzalez. 2023. Gorilla: Large lan-	1508
1452	Swayam Singh, Xiangru Tang, Leandro von Werra,	guage model connected with massive apis. <i>CoRR</i> ,	1509
1453	and Shayne Longpre. 2024. Octopack: Instruction	abs/2305.15334.	1510
1454	tuning code large language models. In <i>The Twelfth</i>		
1455	<i>International Conference on Learning Representa-</i>	Debalina Ghosh Paul, Hong Zhu, and Ian Bayley. 2024.	1511
1456	<i>tions, ICLR 2024, Vienna, Austria, May 7-11, 2024.</i>	Sceneval: A benchmark for scenario-based evalua-	1512
1457	OpenReview.net.	tion of code generation. In <i>IEEE International Con-</i>	1513
		<i>ference on Artificial Intelligence Testing, AITest 2024,</i>	1514
1458	Hoan Anh Nguyen, Tien N. Nguyen, Danny Dig,	<i>Shanghai, China, July 15-18, 2024</i> , pages 55–63.	1515
1459	Son Nguyen, Hieu Tran, and Michael Hilton. 2019.	IEEE.	1516
1460	Graph-based mining of in-the-wild, fine-grained, se-	Nathaniel Ross Pinckney, Christopher Batten, Mingjie	1517
1461	mantic code change patterns. In <i>Proceedings of the</i>	Liu, Haoxing Ren, and Brucek Khailany. 2024.	1518
1462	<i>41st International Conference on Software Engineer-</i>	Revisiting verilogeval: Newer llms, in-context	1519
1463	<i>ing, ICSE 2019, Montreal, QC, Canada, May 25-31,</i>	learning, and specification-to-rtl tasks. <i>CoRR</i> ,	1520
1464	<i>2019</i> , pages 819–830. IEEE / ACM.	abs/2408.11053.	1521

- Disha Shrivastava, Hugo Larochelle, and Daniel Tarlow. 2023b. [Repository-level prompt generation for large language models of code](#). In *International Conference on Machine Learning, ICML 2023, 23-29 July 2023, Honolulu, Hawaii, USA*, volume 202 of *Proceedings of Machine Learning Research*, pages 31693–31715. PMLR.
- Alexander Shypula, Aman Madaan, Yimeng Zeng, Uri Alon, Jacob R. Gardner, Yiming Yang, Milad Hashemi, Graham Neubig, Parthasarathy Ranganathan, Osbert Bastani, and Amir Yazdanbakhsh. 2024. [Learning performance-improving code edits](#). In *The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024*. OpenReview.net.
- Chenglei Si, Yanzhe Zhang, Zhengyuan Yang, Ruibo Liu, and Diyi Yang. 2024. [Design2code: How far are we from automating front-end engineering?](#) *CoRR*, abs/2403.03163.
- Manav Singhal, Tushar Aggarwal, Abhijeet Awasthi, Nagarajan Natarajan, and Aditya Kanade. 2024. [Nofuneval: Funny how code lms falter on requirements beyond functional correctness](#). *CoRR*, abs/2401.15963.
- Patrick Suppes, Joseph L Zinnes, et al. 1962. *Basic measurement theory*.
- Jeffrey Svajlenko, Judith F. Islam, Iman Keivanloo, Chanchal Kumar Roy, and Mohammad Mamun Mia. 2014. [Towards a big data curated benchmark of inter-project code clones](#). In *30th IEEE International Conference on Software Maintenance and Evolution, Victoria, BC, Canada, September 29 - October 3, 2014*, pages 476–480. IEEE Computer Society.
- Xiangru Tang, Bill Qian, Rick Gao, Jiakang Chen, Xinyun Chen, and Mark B Gerstein. 2024. Biocoder: a benchmark for bioinformatics code generation with large language models. *Bioinformatics*, 40(Supplement_1):i266–i276.
- Wei Tao, Yanlin Wang, Ensheng Shi, Lun Du, Shi Han, Hongyu Zhang, Dongmei Zhang, and Wenqiang Zhang. 2022. [A large-scale empirical study of commit message generation: models, datasets and evaluation](#). *Empir. Softw. Eng.*, 27(7):198.
- Runchu Tian, Yining Ye, Yujia Qin, Xin Cong, Yankai Lin, Yinxu Pan, Yesai Wu, Haotian Hui, Weichuan Liu, Zhiyuan Liu, and Maosong Sun. 2024. [Debugbench: Evaluating debugging capability of large language models](#). In *Findings of the Association for Computational Linguistics, ACL 2024, Bangkok, Thailand and virtual meeting, August 11-16, 2024*, pages 4173–4198. Association for Computational Linguistics.
- Michele Tufano, Shao Kun Deng, Neel Sundaresan, and Alexey Svyatkovskiy. 2022. [METHODS2TEST: A dataset of focal methods mapped to test cases](#). In *19th IEEE/ACM International Conference on Mining Software Repositories, MSR 2022, Pittsburgh, PA, USA, May 23-24, 2022*, pages 299–303. ACM.
- Michele Tufano, Cody Watson, Gabriele Bavota, Massimiliano Di Penta, Martin White, and Denys Poshyvanyk. 2019. [An empirical study on learning bug-fixing patches in the wild via neural machine translation](#). *ACM Trans. Softw. Eng. Methodol.*, 28(4):19:1–19:29.
- Prashanth Vijayaraghavan, Luyao Shi, Stefano Ambrogio, Charles Mackin, Apoorva Nitsure, David Beymer, and Ehsan Degan. 2024. [Vhdl-eval: A framework for evaluating large language models in VHDL code generation](#). *CoRR*, abs/2406.04379.
- Yao Wan, Zhou Zhao, Min Yang, Guandong Xu, Haochao Ying, Jian Wu, and Philip S. Yu. 2018. [Improving automatic source code summarization via deep reinforcement learning](#). In *Proceedings of the 33rd ACM/IEEE International Conference on Automated Software Engineering, ASE 2018, Montpellier, France, September 3-7, 2018*, pages 397–407. ACM.
- Junjie Wang, Yuchao Huang, Chunyang Chen, Zhe Liu, Song Wang, and Qing Wang. 2024a. Software testing with large language models: Survey, landscape, and vision. *IEEE Transactions on Software Engineering*.
- Ping Wang, Tian Shi, and Chandan K. Reddy. 2020. [Text-to-sql generation for question answering on electronic medical records](#). In *WWW '20: The Web Conference 2020, Taipei, Taiwan, April 20-24, 2020*, pages 350–361. ACM / IW3C2.
- Shuai Wang, Liang Ding, Li Shen, Yong Luo, Bo Du, and Dacheng Tao. 2024b. [OOP: object-oriented programming evaluation benchmark for large language models](#). In *Findings of the Association for Computational Linguistics, ACL 2024, Bangkok, Thailand and virtual meeting, August 11-16, 2024*, pages 13619–13639. Association for Computational Linguistics.
- Wenhan Wang, Chenyuan Yang, Zhijie Wang, Yuheng Huang, Zhaoyang Chu, Da Song, Lingming Zhang, An Ran Chen, and Lei Ma. 2024c. [TESTEVAL: benchmarking large language models for test case generation](#). *CoRR*, abs/2406.04531.
- Yibo Wang, Ying Wang, Tingwei Zhang, Yue Yu, Shing-Chi Cheung, Hai Yu, and Zhiliang Zhu. 2023a. [Can machine learning pipelines be better configured?](#) In *Proceedings of the 31st ACM Joint European Software Engineering Conference and Symposium on the Foundations of Software Engineering, ESEC/FSE 2023*, page 463–475, New York, NY, USA. Association for Computing Machinery.
- Yu Emma Wang, Gu-Yeon Wei, and David Brooks. 2019. [Benchmarking tpu, gpu, and cpu platforms for deep learning](#). *Preprint*, arXiv:1907.10701.
- Zhiruo Wang, Grace Cuenca, Shuyan Zhou, Frank F. Xu, and Graham Neubig. 2023b. [Mconala: A benchmark for code generation from multiple natural languages](#). In *Findings of the Association for Computational Linguistics: EACL 2023, Dubrovnik, Croatia, May 2-6, 2023*, pages 265–273. Association for Computational Linguistics.

1750	Zhiruo Wang, Shuyan Zhou, Daniel Fried, and Graham Neubig. 2022. Execution-based evaluation for open-domain code generation. <i>arXiv preprint arXiv:2212.10481</i> .	1806
1751		1807
1752		1808
1753		1809
1754	Zora Zhiruo Wang, Akari Asai, Xinyan Velocity Yu, Frank F. Xu, Yiqing Xie, Graham Neubig, and Daniel Fried. 2024d. Coderag-bench: Can retrieval augment code generation? <i>CoRR</i> , abs/2406.14497.	1810
1755		1811
1756		1812
1757		1813
1758	Cody Watson, Michele Tufano, Kevin Moran, Gabriele Bavota, and Denys Poshyvanyk. 2020. On learning meaningful assert statements for unit test cases . In <i>ICSE '20: 42nd International Conference on Software Engineering, Seoul, South Korea, 27 June - 19 July, 2020</i> , pages 1398–1409. ACM.	1814
1759		1815
1760		1816
1761		1817
1762		
1763		
1764	Alexander Wei, Nika Haghtalab, and Jacob Steinhardt. 2024. Jailbroken: How does llm safety training fail? <i>Advances in Neural Information Processing Systems</i> , 36.	1818
1765		1819
1766		1820
1767		1821
1768	Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le, and Denny Zhou. 2022. Chain-of-thought prompting elicits reasoning in large language models. In <i>NeurIPS</i> .	1822
1769		1823
1770		1824
1771		1825
1772		
1773	Jiayi Wei, Greg Durrett, and Isil Dillig. 2023. Typet5: Seq2seq type inference using static analysis . In <i>The Eleventh International Conference on Learning Representations, ICLR 2023, Kigali, Rwanda, May 1-5, 2023</i> . OpenReview.net.	1826
1774		1827
1775		1828
1776		1829
1777		1830
1778	Chengyue Wu, Yixiao Ge, Qiushan Guo, Jiahao Wang, Zhixuan Liang, Zeyu Lu, Ying Shan, and Ping Luo. 2024a. Plot2code: A comprehensive benchmark for evaluating multi-modal large language models in code generation from scientific plots . <i>CoRR</i> , abs/2405.07990.	1831
1779		1832
1780		1833
1781		1834
1782		1835
1783		1836
1784		1837
1785	Tongtong Wu, Weigang Wu, Xingyu Wang, Kang Xu, Suyu Ma, Bo Jiang, Ping Yang, Zhenchang Xing, Yuan-Fang Li, and Gholamreza Haffari. 2024b. Versi-code: Towards version-controllable code generation . <i>CoRR</i> , abs/2406.07411.	1838
1786		1839
1787		1840
1788		1841
1789	Chunqiu Steven Xia, Yinlin Deng, and Lingming Zhang. 2024a. Top leaderboard ranking = top coding proficiency, always? <i>evoeval: Evolving coding benchmarks via LLM</i> . <i>CoRR</i> , abs/2403.19114.	1842
1790		1843
1791		1844
1792		1845
1793		1846
1794	Yinghui Xia, Yuyan Chen, Tianyu Shi, Jun Wang, and Jinsong Yang. 2024b. Aicodereval: Improving AI domain code generation of large language models . <i>CoRR</i> , abs/2406.04712.	1847
1795		
1796		
1797	Jie Xiao, Qianyi Huang, Xu Chen, and Chen Tian. 2024. Large language model performance benchmarking on mobile platforms: A thorough evaluation . <i>Preprint</i> , arXiv:2410.03613.	1848
1798		1849
1799		1850
1800		1851
1801	Ruiyang Xu, Jialun Cao, Yaojie Lu, Hongyu Lin, Xianpei Han, Ben He, Shing-Chi Cheung, and Le Sun. 2024. Cruxeval-x: A benchmark for multilingual code reasoning, understanding and execution . <i>CoRR</i> , abs/2408.13001.	1852
1802		1853
1803		1854
1804		1855
1805		1856
		1857
		1858
		1859
	Ankit Yadav, Himanshu Beniwal, and Mayank Singh. 2024a. Pythonsaga: Redefining the benchmark to evaluate code generating llms . In <i>Findings of the Association for Computational Linguistics: EMNLP 2024</i> , pages 17113–17126.	1860
		1861
		1862
	Ankit Yadav, Himanshu Beniwal, and Mayank Singh. 2024b. Pythonsaga: Redefining the benchmark to evaluate code generating llms . In <i>Findings of the Association for Computational Linguistics: EMNLP 2024, Miami, Florida, USA, November 12-16, 2024</i> , pages 17113–17126. Association for Computational Linguistics.	
	Shuhan Yan, Hang Yu, Yuting Chen, Beijun Shen, and Lingxiao Jiang. 2020. Are the code snippets what we are searching for? A benchmark and an empirical study on code search with natural-language queries . In <i>27th IEEE International Conference on Software Analysis, Evolution and Reengineering, SANER 2020, London, ON, Canada, February 18-21, 2020</i> , pages 344–354. IEEE.	
	Weixiang Yan, Yuchen Tian, Yunzhe Li, Qian Chen, and Wen Wang. 2023. Codetransocean: A comprehensive multilingual benchmark for code translation . In <i>Findings of the Association for Computational Linguistics: EMNLP 2023, Singapore, December 6-10, 2023</i> , pages 5067–5089. Association for Computational Linguistics.	
	Hongyu Yang, Liyang He, Min Hou, Shuanghong Shen, Rui Li, Jiahui Hou, Jianhui Ma, and Junda Zhao. 2024a. Aligning llms through multi-perspective user preference ranking-based feedback for programming question answering . <i>CoRR</i> , abs/2406.00037.	
	Zhiyu Yang, Zihan Zhou, Shuo Wang, Xin Cong, Xu Han, Yukun Yan, Zhenghao Liu, Zhixing Tan, Pengyuan Liu, Dong Yu, Zhiyuan Liu, Xiaodong Shi, and Maosong Sun. 2024b. Matplotagent: Method and evaluation for llm-based agentic scientific data visualization . In <i>Findings of the Association for Computational Linguistics, ACL 2024, Bangkok, Thailand and virtual meeting, August 11-16, 2024</i> , pages 11789–11804. Association for Computational Linguistics.	
	Ziyu Yao, Daniel S. Weld, Wei-Peng Chen, and Huan Sun. 2018. Staqa: A systematically mined question-code dataset from stack overflow . In <i>Proceedings of the 2018 World Wide Web Conference on World Wide Web, WWW 2018, Lyon, France, April 23-27, 2018</i> , pages 1693–1703. ACM.	
	Junjie Ye, Xuanting Chen, Nuo Xu, Can Zu, Zekai Shao, Shichun Liu, Yuhua Cui, Zeyang Zhou, Chao Gong, Yang Shen, Jie Zhou, Siming Chen, Tao Gui, Qi Zhang, and Xuanjing Huang. 2023. A comprehensive capability analysis of gpt-3 and gpt-3.5 series models . <i>Preprint</i> , arXiv:2303.10420.	
	Pengcheng Yin, Bowen Deng, Edgar Chen, Bogdan Vasilescu, and Graham Neubig. 2018. Learning to mine aligned code and natural language pairs from	

1863	stack overflow . In <i>International Conference on Mining Software Repositories</i> , MSR, pages 476–486. ACM.	1921
1864		1922
1865		1923
1866	Pengcheng Yin, Wen-Ding Li, Kefan Xiao, Abhishek Rao, Yeming Wen, Kensen Shi, Joshua Howland, Paige Bailey, Michele Catasta, Henryk Michalewski, Oleksandr Polozov, and Charles Sutton. 2023. Natural language to code generation in interactive data science notebooks . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , ACL 2023, Toronto, Canada, July 9-14, 2023, pages 126–173. Association for Computational Linguistics.	1924
1867		1925
1868		1926
1869		1927
1870		1928
1871		
1872		1929
1873		1930
1874		1931
1875		1932
1876	Hao Yu, Bo Shen, Dezhi Ran, Jiabin Zhang, Qi Zhang, Yuchi Ma, Guangtai Liang, Ying Li, Tao Xie, and Qianxiang Wang. 2023. Codereval: A benchmark of pragmatic code generation with generative pre-trained models. <i>arXiv preprint arXiv:2302.00288</i> .	1933
1877		1934
1878		1935
1879		1936
1880		1937
1881	Tao Yu, Rui Zhang, Heyang Er, Suyi Li, Eric Xue, Bo Pang, Xi Victoria Lin, Yi Chern Tan, Tianze Shi, Zihan Li, Youxuan Jiang, Michihiro Yasunaga, Sungrok Shim, Tao Chen, Alexander R. Fabbri, Zifan Li, Luyao Chen, Yuwen Zhang, Shreya Dixit, Vincent Zhang, Caiming Xiong, Richard Socher, Walter S. Lasecki, and Dragomir R. Radev. 2019a. Cosql: A conversational text-to-sql challenge towards cross-domain natural language interfaces to databases . In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing, EMNLP-IJCNLP 2019, Hong Kong, China, November 3-7, 2019</i> , pages 1962–1979. Association for Computational Linguistics.	1938
1882		1939
1883		1940
1884		1941
1885		1942
1886		
1887		1943
1888		1944
1889		1945
1890		1946
1891		1947
1892		
1893		1948
1894		1949
1895		1950
1896	Tao Yu, Rui Zhang, Kai Yang, Michihiro Yasunaga, Dongxu Wang, Zifan Li, James Ma, Irene Li, Qingning Yao, Shanelle Roman, Zilin Zhang, and Dragomir R. Radev. 2018. Spider: A large-scale human-labeled dataset for complex and cross-domain semantic parsing and text-to-sql task . In <i>Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing, Brussels, Belgium, October 31 - November 4, 2018</i> , pages 3911–3921. Association for Computational Linguistics.	1951
1897		1952
1898		1953
1899		1954
1900		1955
1901		1956
1902		
1903		1957
1904		1958
1905		1959
1906	Tao Yu, Rui Zhang, Michihiro Yasunaga, Yi Chern Tan, Xi Victoria Lin, Suyi Li, Heyang Er, Irene Li, Bo Pang, Tao Chen, Emily Ji, Shreya Dixit, David Proctor, Sungrok Shim, Jonathan Kraft, Vincent Zhang, Caiming Xiong, Richard Socher, and Dragomir R. Radev. 2019b. Sparc: Cross-domain semantic parsing in context . In <i>Proceedings of the 57th Conference of the Association for Computational Linguistics, ACL 2019, Florence, Italy, July 28- August 2, 2019, Volume 1: Long Papers</i> , pages 4511–4523. Association for Computational Linguistics.	1960
1907		1961
1908		1962
1909		1963
1910		1964
1911		
1912		1965
1913		1966
1914		1967
1915		1968
1916		
1917	Xiao Yu, Lei Liu, Xing Hu, Jacky Wai Keung, Jin Liu, and Xin Xia. 2024. Fight fire with fire: How much can we trust chatgpt on source code-related tasks? <i>arXiv preprint arXiv:2405.12641</i> .	1969
1918		1970
1919		1971
1920		1972
		1973
		1974
		1975
		1976
		1977
	Xiaoqing Yu, Tianlong Chen, Zhengjie Yu, Huiyu Li, Yang Yang, Xiaoqian Jiang, and Anxiao Jiang. 2020. Dataset and enhanced model for eligibility criteria-to-sql semantic parsing . In <i>Proceedings of The 12th Language Resources and Evaluation Conference, LREC 2020, Marseille, France, May 11-16, 2020</i> , pages 5829–5837. European Language Resources Association.	
	Yuliang Yuan, Wenxiang Jiao, Wenxuan Wang, Jen tse Huang, Pinjia He, Shuming Shi, and Zhaopeng Tu. 2024a. Gpt-4 is too smart to be safe: Stealthy chat with llms via cipher . <i>Preprint</i> , arXiv:2308.06463.	
	Zhiqiang Yuan, Mingwei Liu, Shiji Ding, Kaixin Wang, Yixuan Chen, Xin Peng, and Yiling Lou. 2024b. Evaluating and improving chatgpt for unit test generation. <i>Proceedings of the ACM on Software Engineering</i> , 1(FSE):1703–1726.	
	Zhiqiang Yuan, Yiling Lou, Mingwei Liu, Shiji Ding, Kaixin Wang, Yixuan Chen, and Xin Peng. 2023a. No more manual tests? evaluating and improving chatgpt for unit test generation. <i>arXiv preprint arXiv:2305.04207</i> .	
	Zhiqiang Yuan, Yiling Lou, Mingwei Liu, Shiji Ding, Kaixin Wang, Yixuan Chen, and Xin Peng. 2023b. No more manual tests? evaluating and improving chatgpt for unit test generation . <i>CoRR</i> , abs/2305.04207.	
	Xiang Yue, Yuansheng Ni, Kai Zhang, Tianyu Zheng, Ruochi Liu, Ge Zhang, Samuel Stevens, Dongfu Jiang, Weiming Ren, Yuxuan Sun, Cong Wei, Botao Yu, Ruibin Yuan, Renliang Sun, Ming Yin, Boyuan Zheng, Zhenzhu Yang, Yibo Liu, Wenhao Huang, Huan Sun, Yu Su, and Wenhao Chen. 2024. Mmmu: A massive multi-discipline multimodal understanding and reasoning benchmark for expert agi . <i>Preprint</i> , arXiv:2311.16502.	
	Sukmin Yun, Haokun Lin, Rusiru Thushara, Mohammad Qazim Bhat, Yongxin Wang, Zutao Jiang, Mingkai Deng, Jinhong Wang, Tianhua Tao, Junbo Li, Haonan Li, Preslav Nakov, Timothy Baldwin, Zhengzhong Liu, Eric P. Xing, Xiaodan Liang, and Zhiqiang Shen. 2024. Web2code: A large-scale webpage-to-code dataset and evaluation framework for multimodal llms . <i>CoRR</i> , abs/2406.20098.	
	Daoguang Zan, Bei Chen, Zeqi Lin, Bei Guan, Yongji Wang, and Jian-Guang Lou. 2022a. When language model meets private library. <i>arXiv preprint arXiv:2210.17236</i> .	
	Daoguang Zan, Bei Chen, Dejian Yang, Zeqi Lin, Minsu Kim, Bei Guan, Yongji Wang, Weizhu Chen, and Jian-Guang Lou. 2022b. Cert: continual pre-training on sketches for library-oriented code generation. <i>arXiv preprint arXiv:2206.06888</i> .	
	Zhengran Zeng, Yidong Wang, Rui Xie, Wei Ye, and Shikun Zhang. 2024. Coderujb: An executable and unified java benchmark for practical programming scenarios . In <i>Proceedings of the 33rd ACM SIGSOFT</i>	

1978	<i>International Symposium on Software Testing and Analysis, ISSTA 2024, Vienna, Austria, September 16-20, 2024</i> , pages 124–136. ACM.	
1979		
1980		
1981	Kechi Zhang, Jia Li, Ge Li, Xianjie Shi, and Zhi Jin.	
1982	2024a. Codeagent: Enhancing code generation with tool-integrated agent systems for real-world repo-level coding challenges . In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , ACL 2024, Bangkok, Thailand, August 11-16, 2024, pages 13643–13658. Association for Computational Linguistics.	
1983		
1984		
1985		
1986		
1987		
1988		
1989		
1990	Quanjun Zhang, Tongke Zhang, Juan Zhai, Chunrong Fang, Bowen Yu, Weisong Sun, and Zhenyu Chen.	
1991	2023a. A critical review of large language model on software engineering: An example from chatgpt and automated program repair . <i>CoRR</i> , abs/2310.08879.	
1992		
1993		
1994		
1995	Shudan Zhang, Hanlin Zhao, Xiao Liu, Qinkai Zheng, Zehan Qi, Xiaotao Gu, Xiaohan Zhang, Yuxiao Dong, and Jie Tang. 2024b. Naturalcodebench: Examining coding performance mismatch on humaneval and natural user prompts . <i>CoRR</i> , abs/2405.04520.	
1996		
1997		
1998		
1999		
2000	Yusen Zhang, Jun Wang, Zhiguo Wang, and Rui Zhang.	
2001	2023b. Xsemplr: Cross-lingual semantic parsing in multiple natural languages and meaning representations . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , ACL 2023, Toronto, Canada, July 9-14, 2023, pages 15918–15947. Association for Computational Linguistics.	
2002		
2003		
2004		
2005		
2006		
2007		
2008	Ziying Zhang, Lizhen Xu, Zhaokun Jiang, Hongkun Hao, and Rui Wang. 2024c. Multiple-choice questions are efficient and robust LLM evaluators . <i>CoRR</i> , abs/2405.11966.	
2009		
2010		
2011		
2012	Dewu Zheng, Yanlin Wang, Ensheng Shi, Ruikai Zhang, Yuchi Ma, Hongyu Zhang, and Zibin Zheng.	
2013	2024. Towards more realistic evaluation of llm-based code generation: an experimental study and beyond . <i>CoRR</i> , abs/2406.06918.	
2014		
2015		
2016		
2017	Qinkai Zheng, Xiao Xia, Xu Zou, Yuxiao Dong, Shan Wang, Yufei Xue, Lei Shen, Zihan Wang, Andi Wang, Yang Li, Teng Su, Zhilin Yang, and Jie Tang. 2023a. Codegeex: A pre-trained model for code generation with multilingual benchmarking on humaneval-x . In <i>Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, KDD '23</i> , page 5673–5684, New York, NY, USA. Association for Computing Machinery.	
2018		
2019		
2020		
2021		
2022		
2023		
2024		
2025		
2026	Yunhui Zheng, Saurabh Pujar, Burn L. Lewis, Luca Buratti, Edward A. Epstein, Bo Yang, Jim Laredo, Alessandro Morari, and Zhong Su. 2021. D2A: A dataset built for ai-based vulnerability detection methods using differential analysis . In <i>43rd IEEE/ACM International Conference on Software Engineering: Software Engineering in Practice, ICSE (SEIP) 2021, Madrid, Spain, May 25-28, 2021</i> , pages 111–120. IEEE.	
2027		
2028		
2029		
2030		
2031		
2032		
2033		
2034		
	Zibin Zheng, Kaiwen Ning, Yanlin Wang, Jingwen Zhang, Dewu Zheng, Mingxi Ye, and Jiachi Chen. 2023b. A survey of large language models for code: Evolution, benchmarking, and future trends . <i>arXiv preprint arXiv:2311.10372</i> .	2035
		2036
		2037
		2038
		2039
	Victor Zhong, Caiming Xiong, and Richard Socher. 2017. Seq2sql: Generating structured queries from natural language using reinforcement learning . <i>CoRR</i> , abs/1709.00103.	2040
		2041
		2042
		2043
	Yaqin Zhou, Shangqing Liu, Jing Kai Siow, Xiaoning Du, and Yang Liu. 2019. Devign: Effective vulnerability identification by learning comprehensive program semantics via graph neural networks . In <i>Advances in Neural Information Processing Systems 32: Annual Conference on Neural Information Processing Systems 2019, NeurIPS 2019, December 8-14, 2019, Vancouver, BC, Canada</i> , pages 10197–10207.	2044
		2045
		2046
		2047
		2048
		2049
		2050
		2051
	Ming Zhu, Aneesh Jain, Karthik Suresh, Roshan Ravindran, Sindhu Tipirneni, and Chandan K. Reddy. 2022. Xlcost: A benchmark dataset for cross-lingual code intelligence . <i>CoRR</i> , abs/2206.08474.	2052
		2053
		2054
		2055
	Qiming Zhu, Jialun Cao, Yaojie Lu, Hongyu Lin, Xianpei Han, Le Sun, and Shing-Chi Cheung. 2024. Domaineval: An auto-constructed benchmark for multi-domain code generation . <i>CoRR</i> , abs/2408.13204.	2056
		2057
		2058
		2059
	Terry Yue Zhuo, Minh Chien Vu, Jenny Chim, Han Hu, Wenhao Yu, Ratnadira Widayarsi, Imam Nur Bani Yusuf, Haolan Zhan, Junda He, Indraneil Paul, Simon Brunner, Chen Gong, Thong Hoang, Armel Randy Zebaze, Xiaoheng Hong, Wen-Ding Li, Jean Kadour, Ming Xu, Zhihan Zhang, Prateek Yadav, Naman Jain, Alex Gu, Zhoujun Cheng, Jiawei Liu, Qian Liu, Zijian Wang, David Lo, Binyuan Hui, Niklas Muennighoff, Daniel Fried, Xiaoning Du, Harm de Vries, and Leandro von Werra. 2024. Bigcodebench: Benchmarking code generation with diverse function calls and complex instructions . <i>CoRR</i> , abs/2406.15877.	2060
		2061
		2062
		2063
		2064
		2065
		2066
		2067
		2068
		2069
		2070
		2071
		2072
	Deqing Zou, Sujuan Wang, Shouhuai Xu, Zhen Li, and Hai Jin. 2020. μvuldeepecker: A deep learning-based system for multiclass vulnerability detection . <i>CoRR</i> , abs/2001.02334.	2073
		2074
		2075
		2076

A Statistics of studied benchmarks

In this section, we conducted a comprehensive and detailed statistical analysis of the 274 benchmarks collected.

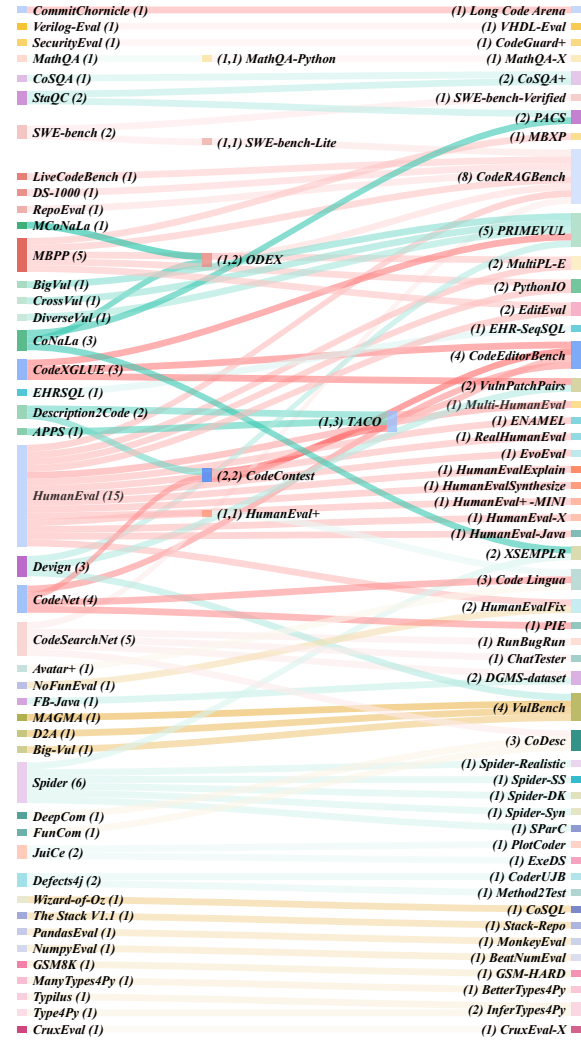


Figure 8: Relationships between Benchmarks

A.1 Profile of Studied Benchmarks

We first show the trend in the development of benchmarks from 2014 to 2024. As shown in Figure 9, the data shows a modest beginning, with only a handful of benchmarks created annually until 2017. From 2018 onwards, there is a noticeable uptrend in benchmark creation, culminating in a significant jump to 149 benchmarks in 2024. This sharp increase indicates a recent heightened interest and demand for comprehensive code-related benchmarks for LLMs, reflecting the evolving complexities and expanding requirements of automated software engineering.

Hierarchy of Benchmarks. Figure 8 visualize

the inheritance relationships among benchmarks, indicating that the benchmarks on the *left serve as sources* for those on the right. It highlights that **18% (50 out of 274) of benchmarks act as data sources**, continuously benefiting the construction of subsequent benchmarks.

Figure 8 reveals that *HumanEval* (Chen et al., 2021a), as the **most significant source** benchmark, benefits at least **15 downstream benchmarks**, followed by MBPP (Austin et al., 2021) and CodeSearchNet (Husain et al., 2019). From the right side of the figure, some benchmarks, like VulBench (Gao et al., 2023b), incorporate methodologies or data from 4 previous benchmarks, and codeRagBench (Wang et al., 2024d) integrates elements from 8 prior benchmarks.

This hierarchical structure among benchmarks also alerts us that the **data quality of a benchmark not only affects its own credibility but can continue to impact others** if it serves as a source. This underscores the importance of adhering to stringent guidelines during benchmark development and highlights the crucial role of **establishing standards** to ensure the integrity and utility of benchmark data across research and development efforts.

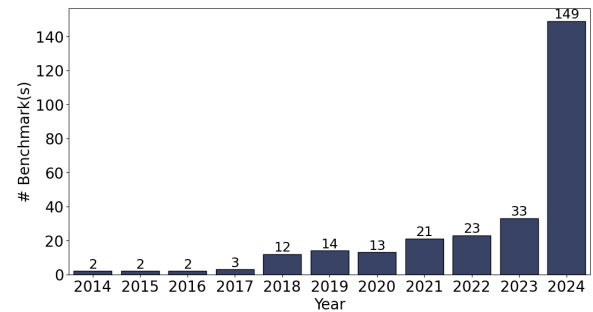


Figure 9: Benchmark Distribution over Years

Annual Trend. Regarding the *coding tasks*, Figure 11 illustrates the distribution of various coding tasks across benchmarks. It is clear that the task of **Code Generation** is most prevalent, with 99 benchmarks focusing on this area according to 36% (99/274) of studied benchmarks, indicating a significant interest in generating code automatically. Program Repair and Defeat Detection are well-represented, with 27 and 25 benchmarks, respectively, reflecting the importance of correcting code and detecting defects.

Citation distribution. We also visualized the citations of 274 code-related benchmarks. The citation statistics were collected on September 1st,

2024. From Figure 10, we can see a clear long-tail trend of the citations, from the highest 2735 (HumanEval (Chen et al., 2021a)) to the lowest 0.

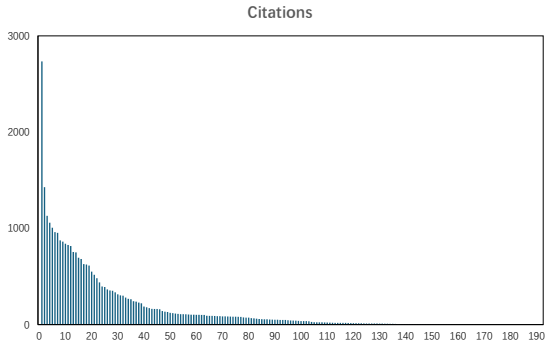


Figure 10: Citation Distribution of Benchmarks

Coding Task. Tasks like Code Summarization and Text2SQL are similarly significant, each with 25 and 22 benchmarks. These tasks focus on making code more understandable and converting natural language queries into SQL queries. Other tasks, such as Code Retrieval, Code Reasoning, and Code Translation, are represented with 18, 17, and 16 benchmarks, respectively. Lesser-represented benchmarks are Test Generation, Code Optimization, and Code Completion, each represented by 8 and 7 benchmarks, indicating the inadequacy of these tasks.

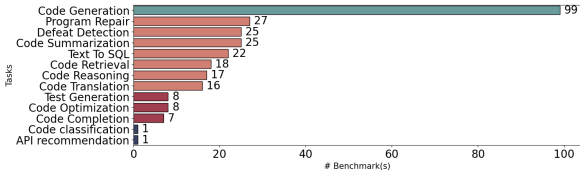


Figure 11: Benchmark Distribution over Tasks

Programming Languages. Figure 12 shows the distribution of benchmarks across various programming languages. The overall trend indicates a strong preference for benchmarking *Python*, which leads with 158 benchmarks, followed by Java and C++, with 107 and 63, respectively. The graph also reveals a diverse range of languages being used. In total, 724 programming languages are studied by these 274 benchmarks. Though some programming languages, such as Kotlin, Swift, and Scala, are less frequently exercised, the benchmarks involving them are tailored to different application needs and technology environments. This distribution shows the existing benchmarks are dominated by three mainstream programming languages, leav-

ing other programming languages less studied and benchmarked.

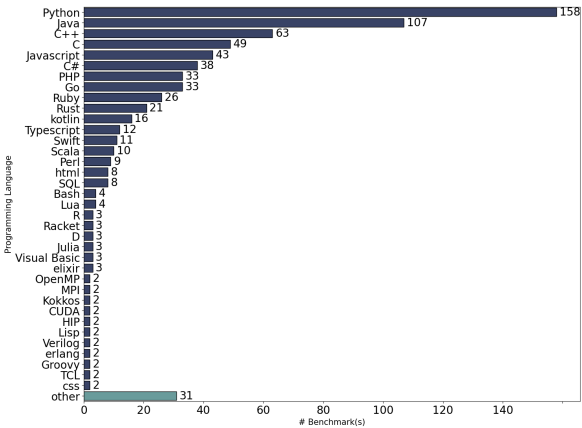


Figure 12: Benchmark Distribution over Programming Language

Natural Language. Figure 13 illustrates the distribution of benchmarks for different natural languages. The bar chart overwhelmingly shows that English is the dominant language, with 192 benchmarks highlighting its ubiquity in research and development. Other languages have significantly fewer benchmarks, with six for Chinese and only two each for Japanese, Russian, and Spanish. The category labeled “Other” includes 20 benchmarks spread across other natural languages, indicating some diversity but limited attention to non-English benchmarks. This distribution highlights the prominence of English in the global research community and also demonstrates the *uneven representation* of natural languages in the studied benchmarks.

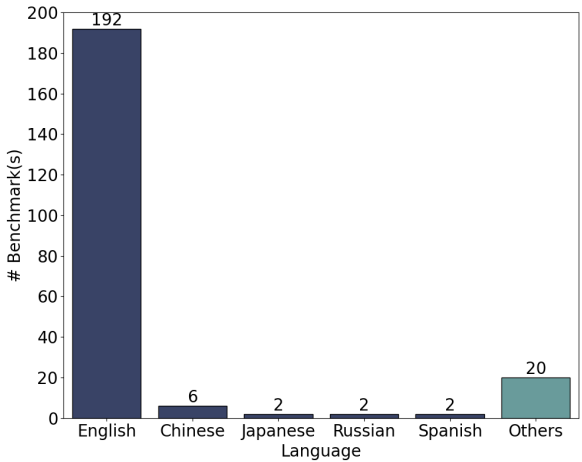


Figure 13: Benchmark Distribution over Natural Language

Modals in the benchmarks. Figure 14 presents

the distribution of benchmarks according to the type of language used in their prompts. The chart shows that the majority, at 47.1%, of the benchmarks use a combination of natural language and programming Language, followed by PL only (31.0%) and NL only (21.9%).

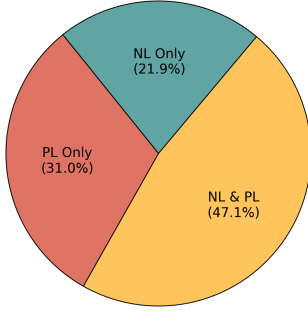


Figure 14: Benchmark Distribution over Modal in Prompt

Granularity. The code snippet in a code-related benchmark varies from statement-level (i.e., one line of code. For example, CoNaLa (Yin et al., 2018) and Math-QA (Amini et al., 2019)), function-level (i.e., a function unit of code. For example, HumanEval (Chen et al., 2021a) and MBPP (Austin et al., 2021)), class-level (i.e., a class with multiple function units of code. For example, ClassEval (Du et al., 2023b)) and project-level (i.e., a project with multiple classes or modules. For example, DevEval (Li et al., 2024a) and JavaBench (Cao et al., 2024a)). Figure 15 illustrates the granularity levels at which benchmarks are typically conducted. The chart shows that the majority of benchmarks, comprising **71.8%, focus on the function level**. Projects constitute 15.0% of the benchmarks. Class-level granularity is the least represented at 2.6%.

The majority of benchmarks are currently at the function level (70+%), followed by the project level (15+%). This indicates that the current major demand is for *assessing individual functions within a single task*, followed by the demand for evaluating functionalities more aligned with *actual project-level code development*.

A.2 Statistics about Benchmark Design

Design of Studied Capabilities. To understand whether benchmark developers recognize the capabilities of LLMs they aim to evaluate, we carefully analyzed 30 representative benchmarks (Appendix C) to see if they clearly specify the capabilities

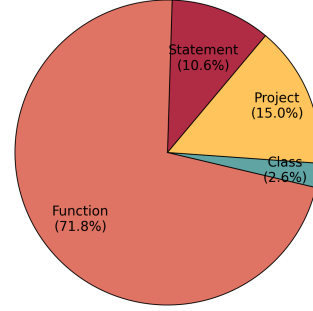


Figure 15: Benchmark Distribution over Granularity

ties being assessed by their benchmarks. As shown in Figure 16, 90% of benchmarks explicitly specify the capabilities (e.g., intention understanding, problem solving, testing, debugging capabilities) to be evaluated, while 10% do not. The statistics show that the most highly cited benchmarks clearly define the assessment capabilities.

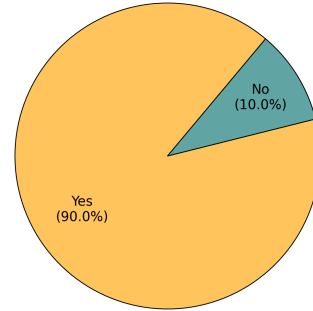


Figure 16: Benchmark Distribution Over Capabilities Consideration

Furthermore, we investigated the 30 focused benchmarks and identified a case (Figure 17) from MBPP (Austin et al., 2021) where the case is likely to fall outside of the targeted capability of the benchmark. In particular, MBPP (Austin et al., 2021) aims to “*measure the ability of these models to synthesize short Python programs from natural language descriptions*” for “*entry-level programmers*”. As we can see from Figure 17, the prompt requires LLMs to “*Write a function to calculate the dogs’ years.*” Simply from this description, an entry-level programmer is unlikely to write a correct program without knowing the conversion equation from dogs’ year to dogs’ age. In other words, this case is more about assessing whether LLMs have acquired this specific knowledge rather than evaluating the most fundamental programming skills.

Design of Studied Application Scenarios. Similarly, to understand whether benchmark developers

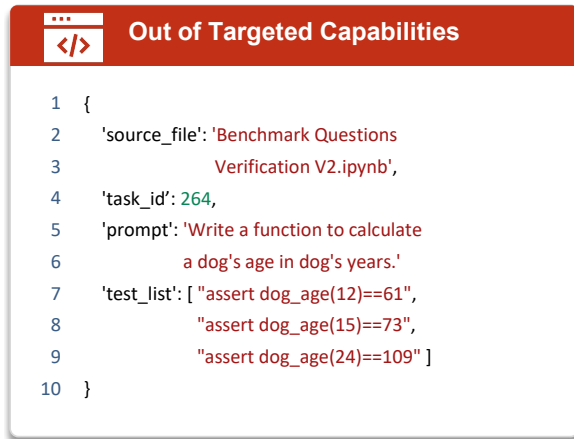


Figure 17: An Example of Out-of-capability Case from MBPP (Austin et al., 2021).

scoped the application scenarios of LLMs they aim to evaluate, we carefully analyzed 30 representative benchmarks (Appendix C) to see whether they explicitly specify the application scenarios their benchmarks target. As shown in Figure 18, 70% representative benchmarks have clearly specified application scenarios (e.g., programming assistant), while the rest do not. Indeed, clearly defining the application scenarios could help benchmark constructors establish precise goals for the design and development of the benchmark, ensuring accuracy in the evaluation.

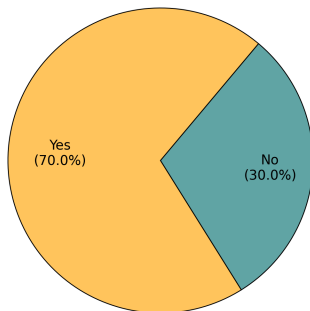


Figure 18: Benchmark Distribution Over Expected Application Scenario Consideration

A.3 Statistics about Data Preparation

A.3.1 Data Preprocessing

Data Deduplication. During benchmark preparation, data cleaning and preprocessing are necessary. However, as shown in Figure 19, only 38% benchmarks have deduplicated the collected data. More than half of them didn’t mention this process.

To investigate the situation, we went through the 30 representative benchmarks (Listed in Ap-

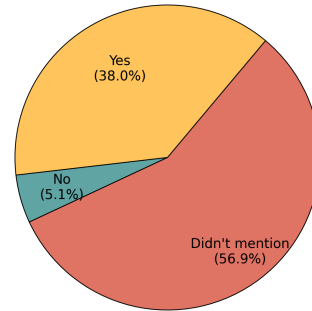


Figure 19: Benchmark Distribution over Deduplication

pendix C) and found two duplicated subjects in MBPP (Austin et al., 2021). Tasks with id 71 and 141 examined the same functionality, i.e., “Write a function to sort a list of elements.”, collected from the same source.

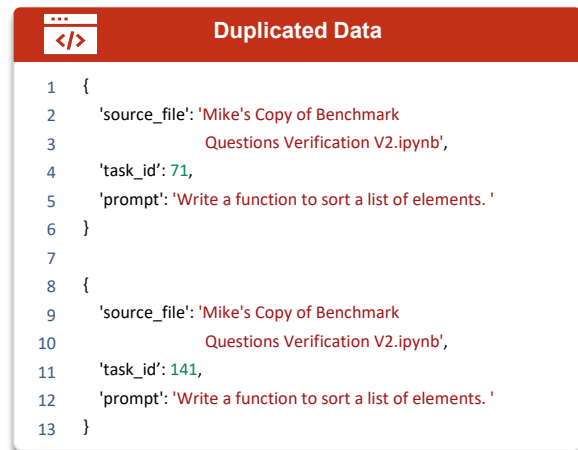


Figure 20: A Counterexample of Rule 16 from MBPP (Austin et al., 2021).

The significance of data preprocessing, such as *deduplication*, is frequently *overlooked* by benchmark builders, leading to data duplication even in highly cited benchmarks.

Data Quality Assurance. Ensuring data quality for the benchmark is essential. However, our statistics (Figure 21) show disappointing results. 67.9% of benchmarks do NOT take any measures for data quality assurance. Among those benchmarks that do incorporate data quality measures, the majority rely on manual checks, which accounts for 22.6%. Other countermeasurements, such as code execution, constitute only 2.2%, while verification using LLMs accounts for 1.5%. Additional methods, such as using download counts as a basis, are also employed.

Additionally, we dived into the 30 representative

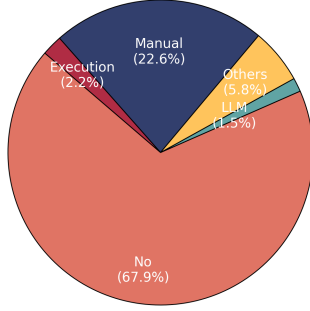


Figure 21: Benchmark Distribution over Quality Assurance Method

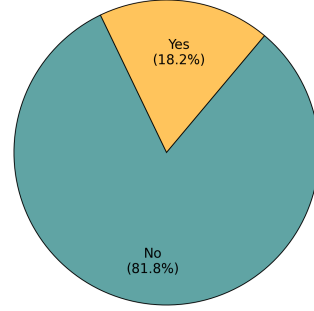


Figure 23: Benchmark Distribution over Quality Assurance on Data Contamination

benchmarks (Listed in Appendix C) and identified an example where the code cannot be executed successfully. As shown in Figure 22, the function `swap()` in line 7 has not been defined, so the execution of the code would fail if the code has been executed. This highlights a significant gap in ensuring the reliability and validity of benchmark data, underscoring the need for more rigorous and automated data quality assurance practices.

```

1 import math
2 def min_Operations (A, B):
3     """ Write a python function to find
4     the minimum operations required
5     to make two numbers equal. """
6     if (A > B):
7         swap(A,B)
8     B = B // math.gcd(A,B);
9     return B - 1

```

Figure 22: An Example from MBPP (Austin et al., 2021) that failed to be executed.

Data Contamination Resolution. Data contamination (Golchin and Surdeanu, 2023; Cao et al., 2024b) threat has been widely discussed. A benchmark with contaminated data may yield overclaimed results, misleading the understanding of the LLMs’ capabilities. According to our statistics (Figure 23 on benchmarks from the year 2023 to 2024 (the duration when most LLMs were launched), most (81.8 %) benchmarks were not aware of and have not taken any measures to alleviate data contamination, being vulnerable to data contamination threat.

A.3.2 Statistics about Data Curation

Ground truth solutions. Figure 24 shows that although the majority (92.3%) of benchmarks pro-

vide reference code as ground truth, there are 5% of benchmarks without reference code. Although it is not compulsory as long as object measurements (e.g., test cases) are provided, **a reference code is still recommended.** Indeed, if a benchmark provides reference code, its reliability tends to be better because it ensures that there are feasible solutions for the tasks involved. This guarantees that the tasks are theoretical and practically solvable, enhancing the benchmark’s usefulness and credibility in real-world applications.

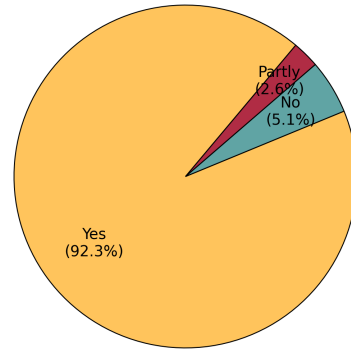


Figure 24: Benchmark Distribution over Solution

Additionally, **the correctness of the ground truth solution** should also be noted. Figure 25 shows an **incorrect** code solution provided in HumanEval (Chen et al., 2021a). This should draw benchmark constructors’ attention to the correctness of the benchmark reference code.

Oracle. An oracle (Barr et al., 2014) is a way to determine whether the output is correct or not. For example, assume the output of LLMs is in the form of code, then an oracle could be running tests against the code and see whether the code can pass all the tests. Figure 26 shows the distribution of types of oracle that are used in these benchmarks. Clear that the **exact match** 41.97% (115/274) and **test case passing** (114/274) 41.6% are the most

```

</>
Incorrect Ground-truth
1 def check_dict_case(dict):
2     """ Given a dictionary, return True if all keys are strings
3     in lower case or all keys are strings in upper case, else
4     return False. The function should return False is the
5     given dictionary is empty. """
6     if len(dict.keys()) == 0:
7         return False
8     else:
9         state = "start"
10        for key in dict.keys():
11            if isinstance(key, str) == False:
12                state = "mixed"
13                break
14            if state == "start":
15                if key.isupper():
16                    state = "upper"
17                elif key.islower():
18                    state = "lower"
19            else:
20                break
21            elif (state == "upper" and not key.isupper())
22                or (state == "lower" and not
23                    key.islower()):
24                state = "mixed"
25                break
26            else:
27                break
28        return state == "upper" or state == "lower"

```

Figure 25: An Example from Humaneval (Chen et al., 2021a) which shows an *incorrect solution* provided in the benchmark.

common oracle used in code-related benchmarks, followed by thresholds (i.e., similarities smaller than a specific threshold).

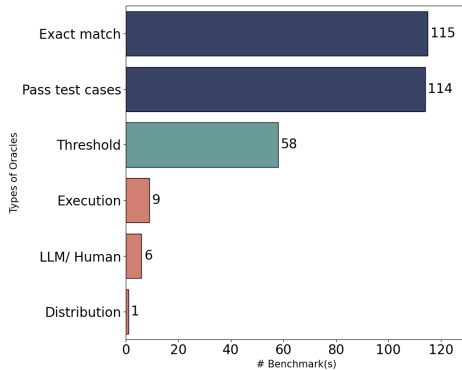


Figure 26: Benchmark Distribution over Test Oracle

Code test coverage (Ivanković et al., 2019), as a common oracle for code-related benchmarks, has been widely adopted to determine the output correctness. It should be considered if a benchmark uses test case passing as a criterion for the correctness of the generated code. Otherwise, a test could be too weak to detect the existence of a defect in the generated code. For example, as pointed out by prior work (Liu et al., 2023a), existing benchmarks such as HumanEval (Chen et al., 2021a) and

MBPP (Austin et al., 2021) still suffer from “insufficient tests”, allowing incorrect code to pass all the tests without capturing the bugs.

Despite its importance, as shown in Figure 27, among the benchmarks that use test cases as the oracle, only 8.7% considered and reported “test coverage” explicitly in their papers, while 87.8% did not mention the test coverage of their provided code.

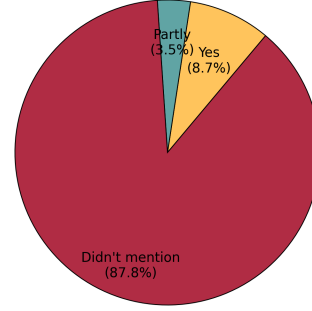


Figure 27: Benchmark Distribution over Test Coverage

Furthermore, we dived into 30 representative benchmarks (Listed in Appendix C) and identified an example (Figure 28) from MBPP (Austin et al., 2021) where the test is incorrect. It alerts us that both the quality of the test and the test adequacy (e.g., code coverage) should be considered.

```

</>
Wrong Example Tests
1 {
2     'source_file': 'charlessutton@: Benchmark
3         Questions Verification V2.ipynb',
4     'task_id': 461,
5     'prompt': 'Write a python function to count the
6         upper case characters in a given string.'
7     'test_list': [ "assert upper_ctr('PYthon') == 1",
8         "assert upper_ctr('BigData') == 1",
9         "assert upper_ctr('program') == 0" ]
10 }

```

Figure 28: An Example of Incorrect Tests from MBPP (Austin et al., 2021).

A.4 Statistics about Evaluation

Studied LLMs. We summarize the number of LLMs that have been evaluated in each benchmark evaluation. Among the 274 benchmarks, 183 of them are evaluated over LLMs, so we show the statistics over them. As shown in Figure 29, over 34% of the benchmarks were evaluated on fewer

than 3 LLMs, with 11.48% benchmarks only evaluated on one LLM. Such evaluation results can hardly be generalized to other LLMs. Furthermore, *more than half of the benchmarks studied fewer than 6 LLMs* ($51\% = (21 + 22 + 20 + 4 + 12 + 15) / 183$).

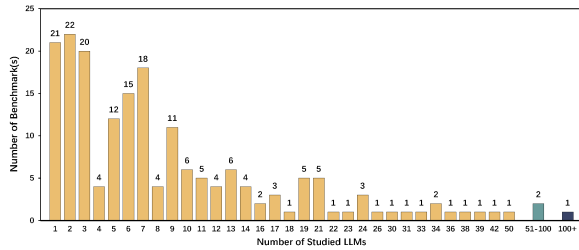


Figure 29: Benchmark Distribution over LLM Experimented

Additionally, we listed the top-10 LLMs by the number of code-related benchmarks they have been evaluated, as shown in Figure 30. GPT series leads significantly with 116 benchmarks, suggesting its widespread adoption and possibly its versatility or superior performance in handling code-related tasks. The rest, including CodeLlama, StarCoder, CodeGen, and others, show varying degrees of involvement, with numbers ranging from 60 down to 24 benchmarks for Claude. This figure may provide a reference for choosing which model to evaluate. In addition, it is worth mentioning that different LLMs should be considered considering different coding tasks.

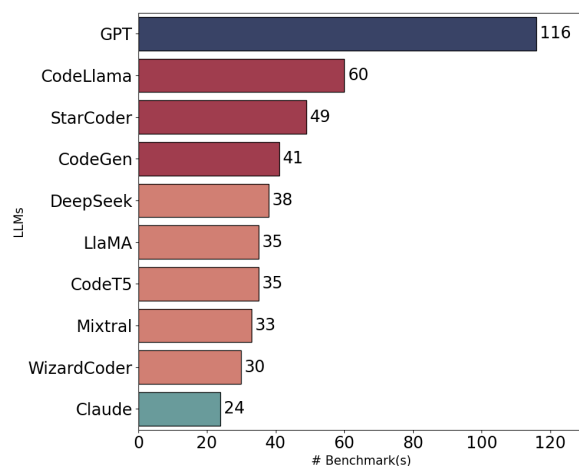


Figure 30: Top-10 Studied LLMs for Code-related Benchmarks

Experiment Environments. The experimental environment (such as the operating system and hardware) is important for the reproduction of the

experiment. However, Figure 32 and Figure 31 highlight a significant gap. A mere 27.4% of benchmarks document the devices used in their experiments, leaving a substantial 72.6% that do not. The situation appears even more dire when considering os, with only 3.6% of benchmarks documenting the OS used, while a staggering 96.4% neglect to record this information.

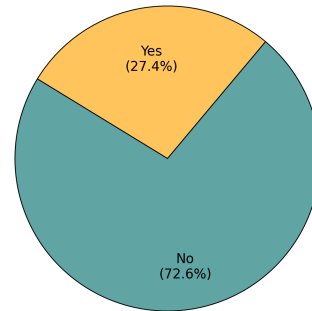


Figure 31: Benchmark Distribution over Recording Experiment Devices

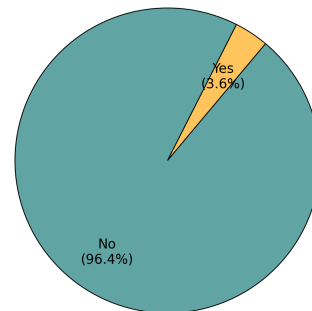


Figure 32: Benchmark Distribution over Recording Experiment OS

Prompting and Prompting Strategies Prompting has a direct impact on the quality of the LLMs' output results (Wei et al., 2022; He et al., 2024a; Jin et al., 2024). So, we summarized whether different prompting strategies have been evaluated and statistics the distribution. Figure 33 shows the usage of four kinds of prompts: zero-shot, few-shot, chain-of-thought, and retrievals (RAG). From Figure 33, we can see that a vast majority (94.9%) benchmarks were evaluated in a zero-shot context setting, while only 21.2% benchmarks were evaluated in a few-shot manner. Even fewer benchmarks were evaluated under the COT and RAG settings, utilized by only 8.8% and 2.6%.

Prompt Quality The prompt quality also greatly impacts the LLM evaluation (He et al., 2024b). So, carefully designing a prompt needs consideration. However, as shown in Figure 34, 73.3% represen-

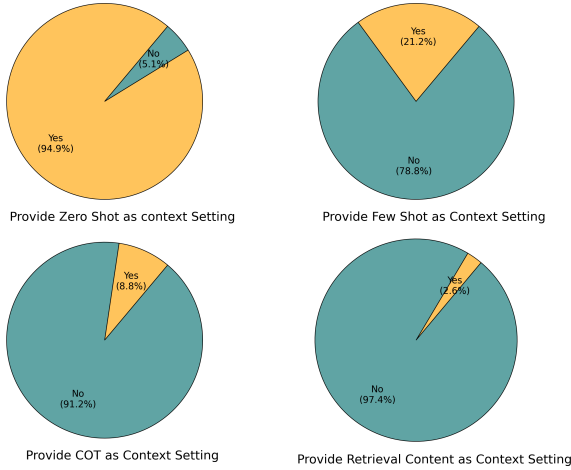


Figure 33: Benchmark Distribution over Context Setting

tative benchmarks (Appendix C) do not validate whether the prompt they used is well-designed.

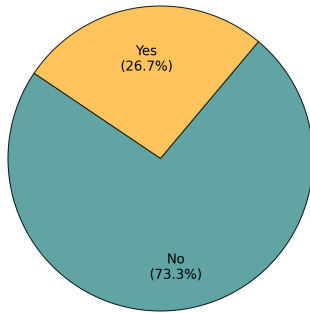


Figure 34: Benchmark Distribution Over Validation of Prompts

Repeated Experiment Given the random nature of LLMs, the experiments are expected to repeat, ensuring the stability and reliability of the results. However, as shown in Figure 35, **only 35.4% benchmarks went through a repeated experiment**, while a majority of 64.6% opted against repeating the experiment. This reflects a lack of awareness regarding the stability and reproducibility of evaluations.

A.5 Statistics about Analysis

Experiment Explanation. Explaining experiment results is crucial for other practitioners to understand what the outcomes mean in the context of the research questions. So, we investigate whether the representative benchmarks (Appendix C) have explained the experiment results. As shown in Figure 36, 70% benchmarks have detailed explanations and analyses of their evaluation results, while still 30% have not. Indeed, an explanation con-

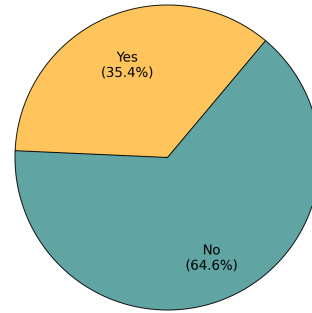


Figure 35: Benchmark Distribution over Repeating the Experiment

tributes to the body of knowledge by making it possible to understand and compare results with previous studies, promoting transparency within the community.

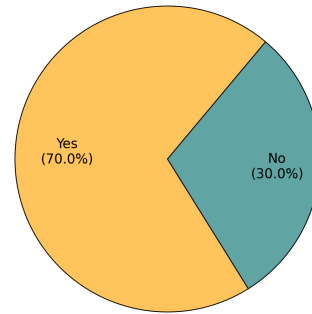


Figure 36: Benchmark Distribution Over Explaining the Experiment

A clear and precise presentation of experimental results is important for enabling robust interpretation and comparison across benchmarks. However, further examination of the 30 representative benchmarks (listed in Appendix C) revealed notable deficiencies in result visualization. As shown in Figure 37, CruxEval (Gu et al., 2024) exhibits unclear experimental result presentation. Specifically, the scatter plot suffers from ambiguous labeling, poor readability of axis values, and inconsistent marker encoding, making it difficult for researchers to extract meaningful insights. Such presentation shortcomings obscure the performance relationships between methods and compromise the benchmark’s usability for fair evaluation. To address these issues, benchmarks should adopt standardized and well-documented visualization practices, ensuring results are interpretable, accessible, and reproducible.

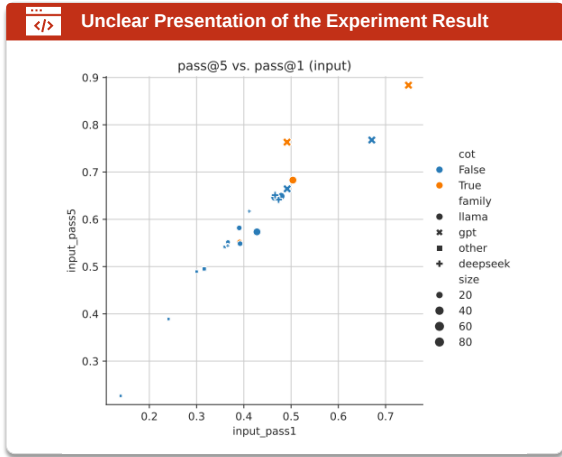


Figure 37: An Example of Unclear Experiment Analysis and Display from CruxEval (Gu et al., 2024)

A.6 Statistics about Release

Data Accessibility. The fundamental requirement for releasing a benchmark is that it must be open-sourced. However, surprisingly, as shown in Figure 38, we observed that 5.1% of the benchmarks are only partially open-sourced (e.g., missing some subjects or tests), and 5.8% are *not open-sourced at all* (e.g., links/web pages are no longer active).

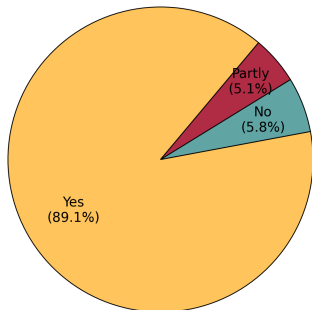


Figure 38: Benchmark Data Availability

Prompt Accessibility. Detailed prompts are essential for ensuring the reproducibility and transparency of code-related benchmarks. However, as shown in Figure 39, we found that 52.6% of benchmarks *do not provide detailed prompts*, limiting the ability to accurately replicate and evaluate the performance of LLMs. This lack of prompt disclosure highlights a gap in benchmark design practices, as prompts are often indispensable for understanding model performance under specific conditions. While 47.4% of benchmarks include such prompts, the absence of comprehensive prompt documentation in over half of the cases raises concerns about the consistency and reproducibility of reported re-

sults.

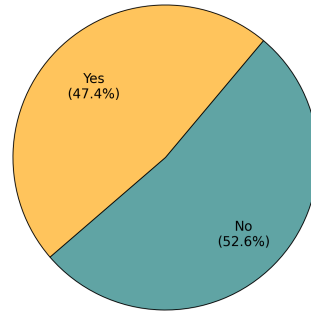


Figure 39: Availability of Prompts

Logging Info Accessibility. Providing detailed logging information, including comprehensive experimental results, is essential for ensuring transparency, verifiability, and reproducibility in benchmarking research. However, as shown in Figure 40, only 16.7% of the benchmarks *make their experimental results publicly available*, while 80.0% fail to disclose this critical information. Alarming, an additional 3.3% provide only partial logging details, further complicating result verification. The absence of complete logging information creates significant barriers for researchers attempting to reproduce experiments or validate reported findings, thereby undermining the reliability of benchmarks. To address this, we emphasize the necessity of making detailed logging information, including intermediate results and metrics, publicly accessible to uphold rigorous scientific standards and foster trustworthy comparisons across models.

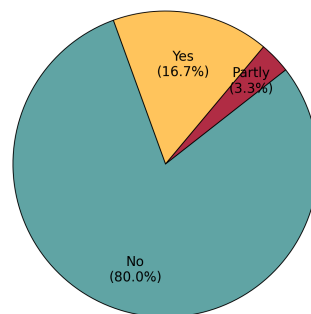


Figure 40: Availability of Logging Information

User Manual Accessibility. A high-quality user manual, such as a well-documented README file, is crucial for enhancing benchmark usability, enabling users to understand the dataset, execute provided scripts, and reproduce results efficiently. However, our analysis revealed that a significant number of benchmarks lack comprehensive user manuals, hindering accessibility and adop-

tion. As depicted in Figure 41, poorly structured or incomplete manuals often omit essential components such as benchmark introductions, usage instructions, and evaluation scripts. This creates unnecessary barriers for researchers who rely on these manuals for setup and experimentation. To address this, we advocate for benchmarks to include clear, standardized user manuals that provide an overview of the benchmark, step-by-step execution guides, and troubleshooting instructions, ensuring a seamless and reproducible user experience.

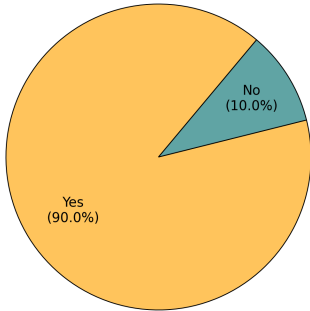


Figure 41: Availability of User Manual

Convenient Evaluation Interface Availability. Providing convenient evaluation interfaces is essential for enhancing the usability and accessibility of benchmarks, enabling researchers to easily reproduce results and compare models. As shown in Figure 42, **16.7% of benchmarks fail to offer any evaluation interfaces**, imposing significant barriers to usability. While a majority of benchmarks (83.3%) provide such interfaces, including command-line tools, Docker images, or scripts, the absence of standardized and user-friendly evaluation tools in a notable minority of cases highlights an area for improvement. Benchmarks without convenient evaluation interfaces require users to spend additional effort in setup and result verification, which can discourage adoption and hinder reproducibility. To address this, we emphasize the importance of releasing benchmarks with well-documented, ready-to-use evaluation pipelines to promote efficient, reliable, and fair benchmarking practices.

Temperature Records. One critical parameter for benchmarking is the temperature setting, which influences stochasticity in LLMs. As shown in Figure 43, we observed that **57.3% of benchmarks fail to record the temperature setting**, hindering reproducibility and fair evaluation. While 42.7% of benchmarks do document this parameter, the majority omission highlights an overlooked yet essential

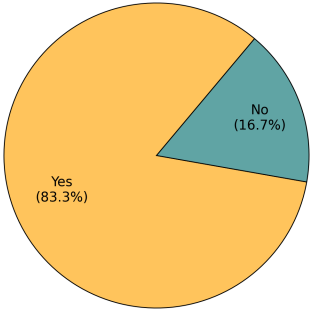


Figure 42: Availability of Convenient Evaluation Interfaces

aspect of benchmark transparency.

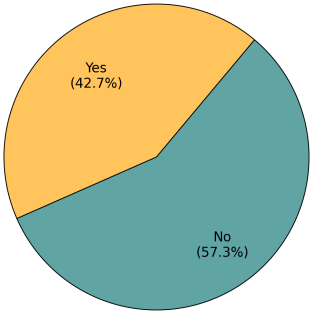


Figure 43: Benchmark Distribution over Recording Temperature

License Provision. Releasing benchmarks under a clear and accessible license is fundamental for legal compliance and ensuring open collaboration. Figure 44 reveals that **19.3% of benchmarks do not provide a license**, limiting their usability and distribution. Encouragingly, 80.7% of benchmarks do include a license, but the lack of licensing in nearly one-fifth of the benchmarks raises concerns about widespread adoption and usage. These findings emphasize the need for standardized practices in benchmark releases to promote legal clarity and accessibility.

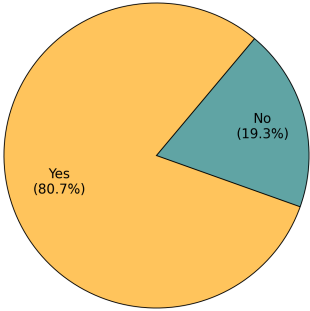


Figure 44: Provision of License

Data Security. Ensuring data security is a critical yet often overlooked aspect of benchmark de-

```
API Key Leakage

1  #!/bin/bash
2
3  # conda activate skg
4  #export WANDB_API_KEY=*****
5  export WANDB_PROJECT=mt5-large_mgeoquery_few-shot
6  export CUDA_LAUNCH_BLOCKING=1
```

CrossVul (Nikitopoulos et al., 2021), where *personal names and email addresses* were unintentionally disclosed. Such leakage poses risks of unauthorized access and resource misuse, potentially compromising systems and research integrity.

```

10942 Individual \fIreadline\FP initialization file
10943 .PD
10944 .SH AUTHORS
10945 Brian Fox, Free Software Foundation
10946 .br
10947 bfox@gnu.org
10948 .PP
10949 Chet Ramey, Case Western Reserve University
10950 .br
10951 chet.ramey@case.edu

```

Usability. Clear and comprehensive documentation is crucial for ensuring the usability of benchmarks, as poorly written instructions can significantly hinder adoption and reproducibility. We dived into the 30 representative benchmarks (listed in Appendix C) and identified an example where the README file provided insufficient and unclear information. As shown in Figure 47, VulDeePecker (Li et al., 2018b) includes a less-than-ideal

Convenient Evaluation Interfaces

First generate the code outputs

```
python3 generate_gpt_codes.py --save /path/to/save_dir
```

Second evaluate the accuracy of the outputted code

```
python3 test_one_solution.py --save /path/to/save_dir
# because the above may fail on account of poorly generated python programs
# we suggest to run a for loop for each problem index against the "all_codes.json"
python3 test_one_solution.py --do
python3 test_one_solution.py --save /path/to/save_dir -i $i ;
done
```

The above will output the accuracy but to run it again once the evaluations have completed execute the line below:

```
python3 test_one_solution.py --save /path/to/save_dir --print_results
```

Third evaluate the bleu scores of the outputted code

```
python3 eval_bleu.py --save /path/to/save_dir
```

B Details of Human Study

The selection of interviewees is pivotal to ensuring the representativeness and relevance of the data collected. This involves identifying individuals

with the knowledge or experience pertinent to the research theme.

To this end, we chose graduate students from SE or AI fields who have published at least one paper. This criterion ensures that participants have research experience and judgment capabilities. The focus on SE and AI fields is due to their likely interest in code benchmarks. Particularly, we aimed to recruit individuals who have published papers on code benchmarks to obtain firsthand feedback from experienced benchmark developers.

B.2 Survey Question Design

Questions. The body of the survey was divided into five stages of benchmark development (following Figure 1), with necessary background information provided for each stage. Each criterion in HOW2BENCH was slightly modified to be in the first-person perspective, making it easier for interviewees to empathize and answer the questions from their own viewpoint. Finally, to facilitate comprehension, questions and instructions were translated into both English and Chinese.

Question Setting. To minimize the effort required from respondents, we designed *single-choice questions* with four options:

- ☐ I found it **important**, and I **have done** it.
- ☐ I found it **important**, although I **haven’t done** it.
- ☐ I found it **not important**, but I **have done** it.
- ☐ I found it **not important**, and I **wouldn’t** do it.

This format is intended to orthogonally explore the correlation between *awareness* and *behavior*.

B.3 Interview Process

Questionnaire Distribution The questionnaire was distributed via online platforms, targeting academic and professional networks related to SE and AI. *The distribution started on October 27, 2024, and ended on November 27th, 2024, lasting one month.*

Results Collection The responses were automatically collected through the online platform used for distribution.

Survey Screening Since the requirement was for participants who have published papers, responses from those selecting “No” to having published a paper were excluded. Also, incomplete surveys where not all questions were answered were also considered invalid and excluded from the analysis.

B.4 Interview Result Analysis

In total, we collected 50 responses. The respondents were from seven regions, including the United States, the United Kingdom, Germany, Australia, China, and others, as shown in Figure 49. Only one survey was invalid due to the respondent selecting “have not published a paper”, leaving **49 valid surveys** for analysis. A breakdown of the respondents’ demographics is shown in Figure 50. The detailed responses for all 55 criteria in HOW2BENCH are shown in Figure 51 and Figure 52.

Geographical distribution of Interviewees

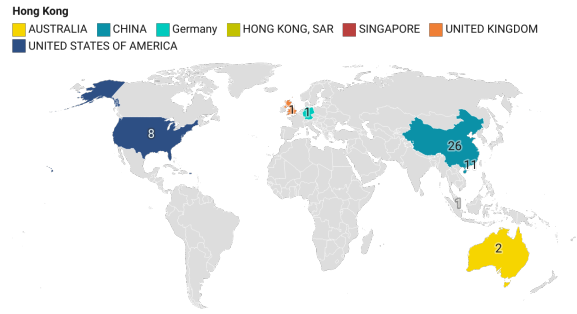


Figure 49: Geographical Distribution of Interviewees

C List of Studied Benchmarks (Focused Ones)

Code Generation: Five with top citations:

- HumanEval (Chen et al., 2021a)
- MBPP (Austin et al., 2021)
- CodeContest (Li et al., 2022)
- leetcodehardgym (Shinn et al., 2023)
- APPS (Hendrycks et al., 2021)

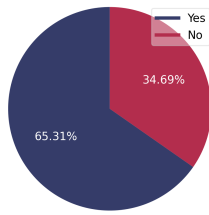
The latest one as of 31/8/2024:

- VerilogEval (Pinckney et al., 2024)

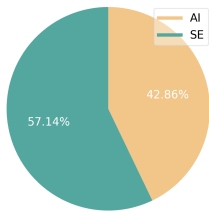
Defect Detection: Five with top citations:

- VulDeePecker (Li et al., 2018b)
- Devign (Zhou et al., 2019)
- Chromium and Debian (Chakraborty et al., 2022)
- μ VulDeePecker (Zou et al., 2020)
- Synthetic Dataset (Hellendoorn et al., 2020)

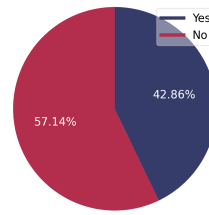
Do you already have a PhD degree or are you a PhD candidate



Do you think you belong more in the SE or AI field?



Have you published at least one paper on the construction of code-related benchmarks (including arXiv)?



Do you think having a checklist for benchmark construction would contribute to the quality of the test set?

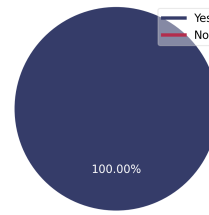


Figure 50: Demography of Interviewees

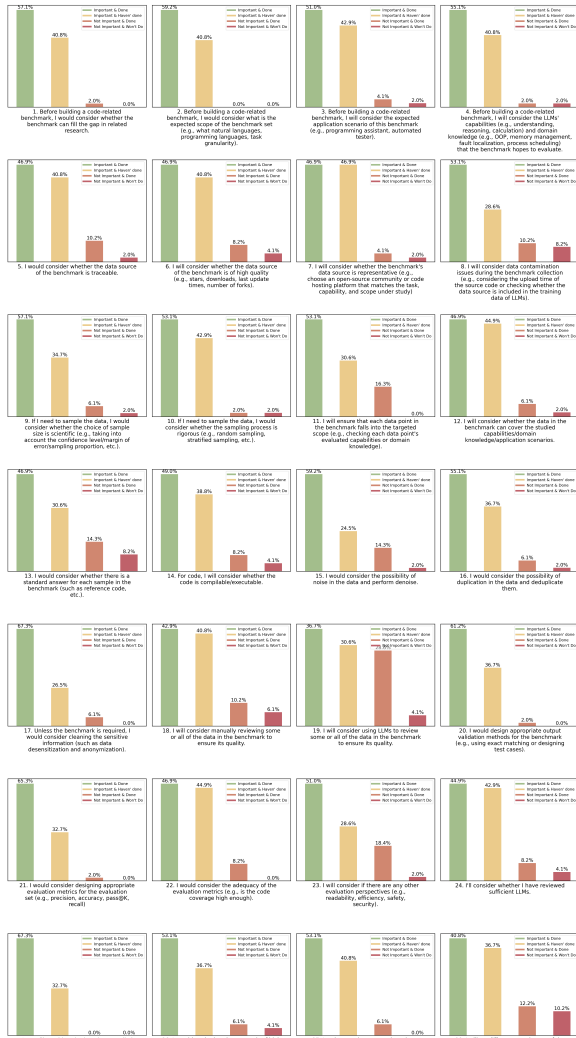


Figure 51: Results of Human Study (Questions 1 - 28)

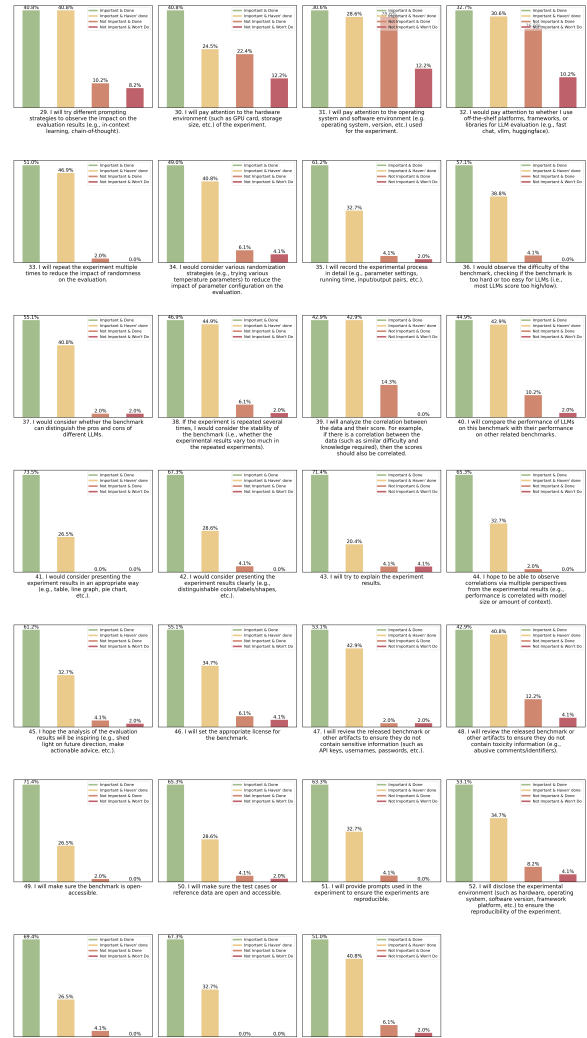


Figure 52: Results of Human Study (Questions 29 - 55)

The latest one as of 31/8/2024:

- VulDetectBench (Liu et al., 2024c)

Program Repair: Five with top citations:

- Defects4J (Just et al., 2014)
- BFP (Tufano et al., 2019)

- MANYBUGS, INTROCLASS (Le Goues et al., 2015)

- HumanEval-Java (Jiang et al., 2023)

- QuixBugs (Prenner et al., 2022)

The latest one as of 31/8/2024:

- DebugBench (Tian et al., 2024)

2691	Code Summarization: Five with top citations:	• VHDL-Eval (Vijayaraghavan et al., 2024)	2724
2692	• CODE-NN (Iyer et al., 2016)	• NaturalCodeBench (Zhang et al., 2024b)	2725
2693	• Java-small/med/large (Alon et al., 2019)	• CodeGuard+ (Fu et al., 2024)	2726
2694	• code-summarization-public (Wan et al., 2018)	• PECC (Haller et al., 2024)	2727
2695	• HumanEvalPack (Muennighoff et al., 2024)	• USACO (Shi et al., 2024b)	2728
2696	• Shrivastava et al. (Shrivastava et al., 2023b)	• ParEval (Nichols et al., 2024)	2729
2697	The latest one as of 31/8/2024:	• MxEval (Athiwaratkun et al., 2022)	2730
2698	• Long Code Arena (Bogomolov et al., 2024)	• MMCode (Li et al., 2024c)	2731
2699	Text To SQL: Five with top citations:	• Plot2Code (Wu et al., 2024a)	2732
2700	• WikiSQL (Zhong et al., 2017)	• ChartMimic (Shi et al., 2024a)	2733
2701	• Spider (Yu et al., 2018)	• DebugBench (Tian et al., 2024)	2734
2702	• Advising (Finegan-Dollak et al., 2018)	• PythonIO (Zhang et al., 2024c)	2735
2703	• BIRD (Li et al., 2023a)	• StaCCQA (Yang et al., 2024a)	2736
2704	• Spider-Realistic (Deng et al., 2021)	• RepoQA (Liu et al., 2024a)	2737
2705	The latest one as of 31/8/2024:	• PRIMEVUL (Ding et al., 2024a)	2738
2706	• AMBROSIA (Saparina and Lapata, 2024)	• VulDetectBench (Liu et al., 2024c)	2739
2707	D List of Studied Benchmarks (Full)	• ProCQA (Li et al., 2024e)	2740
2708	We collected and studied 274 code-related bench-	• CoSQA+ (Gong et al., 2024)	2741
2709	marks. We then listed and grouped them by year.	• JavaBench (Cao et al., 2024a)	2742
2710	2024:	• HumanEvo (Zheng et al., 2024)	2743
2711	• CodeEditorBench (Guo et al., 2024)	• REPOEXEC (Hai et al., 2024)	2744
2712	• MHPP (Dai et al., 2024)	• EHR-SeqSQL (Ryu et al., 2024)	2745
2713	• LiveCodeBench (Jain et al., 2024)	• BookSQL (Kumar et al., 2024)	2746
2714	• CodeAgentBench (Zhang et al., 2024a)	• AMBROSIA (Saparina and Lapata, 2024)	2747
2715	• CruxEval (Gu et al., 2024)	• WUB, WCGB (Yun et al., 2024)	2748
2716	• BigCodeBench (Zhuo et al., 2024)	• RES-Q (LaBash et al., 2024)	2749
2717	• OOPEval (Wang et al., 2024b)	• PythonSaga (Yadav et al., 2024b)	2750
2718	• DevEval (Li et al., 2024a)	• Mercury (Du et al., 2024)	2751
2719	• Long Code Arena (Bogomolov et al., 2024)	• ENAMEL (Qiu et al., 2024b)	2752
2720	• CodeRAGBench (Wang et al., 2024d)	• RealHumanEval (Mozannar et al., 2024)	2753
2721	• ScenEval (Paul et al., 2024)	• CoderUJB (Zeng et al., 2024)	2754
2722	• AICoderEval (Xia et al., 2024b)	• EvoEval (Xia et al., 2024a)	2755
2723	• VersiCode (Wu et al., 2024b)		

2756	• ML-Bench (Liu et al., 2023c)	• EHRSQL (Lee et al., 2023)	2788
2757	• VerilogEval (Pinckney et al., 2024)	• Spider2-V (Cao et al., 2024c)	2789
2758	• CodeApex (Fu et al., 2023)	• TESTEVAL (Wang et al., 2024c)	2790
2759	• HumanEvalPack (Muennighoff et al., 2024)	• ChatTester (Yuan et al., 2023b)	2791
2760	• HumanEval+ (Liu et al., 2023b)	• Code Lingua (Pan et al., 2024)	2792
2761	• HumanEval-X (Zheng et al., 2023a)	• EffiBench (HUANG et al., 2024)	2793
2762	• XCodeEval (Khan et al., 2024)	• CRUXEval-X (Xu et al., 2024)	2794
2763	• CoderEval (Yu et al., 2023)	• DomainEval (Zhu et al., 2024)	2795
2764	• CodeXGLUE (Lu et al., 2021)	2023:	2796
2765	• VulnPatchPairs (Risse and Böhme, 2024)	• MCoNaLa (Wang et al., 2023b)	2797
2766	• WikiSQL (Zhong et al., 2017)	• MultiPL-E (Cassano et al., 2022)	2798
2767	• CrossCodeEval (Ding et al., 2023)	• ODEX (Wang et al., 2022)	2799
2768	• SWE-bench (Jimenez et al., 2024)	• TACO (Li et al., 2023b)	2800
2769	• BAIRI et al. (Bairi et al., 2024)	• DOTPROMPTS, MGDMI- CROBENCH (Agrawal et al., 2023)	2801 2802
2770	• BioCoder (Tang et al., 2024)	• StudentEval (Babe et al., 2024)	2803
2771	• RepoBench (Liu et al., 2024b)	• CodeTransOcean (Yan et al., 2023)	2804
2772	• NoFunEval (Singhal et al., 2024)	• G-TransEval (Jiao et al., 2023)	2805
2773	• CoCoMIC (Ding et al., 2024b)	• AVATAR (Ahmad et al., 2023)	2806
2774	• Java-small/med/large (Alon et al., 2019)	• RunBugRun (Prenner and Robbes, 2023)	2807
2775	• FixEval (Haque et al., 2023)	• VulBench (Gao et al., 2023b)	2808
2776	• CommitBench (Schall et al., 2024)	• DiverseVul (Chen et al., 2023)	2809
2777	• InfiAgent-DABench (Hu et al., 2024)	• Hellendoorn et al. (Hellendoorn et al., 2020)	2810
2778	• InfiBench (Li et al., 2024d)	• XSemPLR (Zhang et al., 2023b)	2811
2779	• Design2Code (Si et al., 2024)	• BIRD (Li et al., 2023a)	2812
2780	• MatPlotBench (Yang et al., 2024b)	• Stack-Repo (Shrivastava et al., 2023a)	2813
2781	• EditEval (Li et al., 2024b)	• RepoEval (Liao et al., 2024)	2814
2782	• D1, D2, D3 (Huang et al., 2024)	• MTPB (Nijkamp et al., 2022)	2815
2783	• RepoEval (Liao et al., 2024)	• ARCADE (Yin et al., 2023)	2816
2784	• BetterTypes4Py, InferTypes4Py (Wei et al., 2023)	• Shrivastava et al. (Shrivastava et al., 2023b)	2817
2785	• HumanEval-Java (Jiang et al., 2023)	• Grag et al. (Garg et al., 2022)	2818
2786	• PIE (Shypula et al., 2024)	• GSM-HARD (Gao et al., 2023a)	2819
2787	• EvalGPTFix (Zhang et al., 2023a)		

2820	• InferredBugs (Jin et al., 2023)	• Berabi et al. (Berabi et al., 2021)	2852
2821	• LeetcodeHardGym (Shinn et al., 2023)	• CrossVul (Nikitopoulos et al., 2021)	2853
2822	• APIBench (Patil et al., 2023)	• PYPIBUGS, RANDOMBUGS (Allamanis et al., 2021)	2854 2855
2823	• ClassEval (Du et al., 2023b)	• D2A (Zheng et al., 2021)	2856
2824	• CommitChronicle (Eliseeva et al., 2023)	• CodeQA (Liu and Wan, 2021)	2857
2825	• TeCo (Nie et al., 2023)	• Spider-DK (Gan et al., 2021b)	2858
2826	• TESTPILOT (Schäfer et al., 2024)	• KaggleDBQA (Lee et al., 2021)	2859
2827	2022:	• SEDE (Hazoom et al., 2021)	2860
2828	• AixBench (Hao et al., 2022)	• Spider-Syn (Gan et al., 2021a)	2861
2829	• TypeBugs (Oh and Oh, 2022)	• CoDesc (Hasan et al., 2021)	2862
2830	• XLCoST (Zhu et al., 2022)	• Methods2Test (Tufano et al., 2022)	2863
2831	• CS1QA (Lee et al., 2022)	• Rozière et al. (Rozière et al., 2022)	2864
2832	• Chromium and Debian (Chakraborty et al., 2022)	2020:	2865
2833	• Spider-Realistic (Deng et al., 2021)	• Lachaux&Roziere et al. (Rozière et al., 2020)	2866
2834	• Spider-SS (Gan et al., 2022)	• μ VulDeePecker (Zou et al., 2020)	2867
2835	• DSP (Chandel et al., 2022)	• CosBench (Yan et al., 2020)	2868
2836	• CodeContest (Li et al., 2022)	• PACS (Heyman and Cutsem, 2020)	2869
2837	• PandasEval, NumpyEval (Zan et al., 2022b)	• Criteria2SQL (Yu et al., 2020)	2870
2838	• TorchDataEval, MonkeyEval, BeatNu-	• SQUALL (Shi et al., 2020)	2871
2839	mEval (Zan et al., 2022a)	• Hu et al. (Hu et al., 2019)	2872
2840	• DS-1000 (Lai et al., 2023)	• CodeSearchNet Challenge (Husain et al., 2019)	2873
2841	• MCMD (Tao et al., 2022)	• MIMICSQL (Wang et al., 2020)	2874
2842	• ExeDS (Huang et al., 2022)	• Atlas (Watson et al., 2020)	2875
2843	• QuixBugs (Prenner et al., 2022)	• Liu et al. (Liu et al., 2022)	2876
2844	• ManyTypes4Py v0.7 (Mir et al., 2022)	• Android (Agarwal et al., 2020)	2877
2845	2021:	• CCSD (Liu et al., 2021)	2878
2846	• SySeVR (Li et al., 2018a)	2019:	2879
2847	• Ling&Wu et al. (Ling et al., 2021)	• BFP (Tufano et al., 2019)	2880
2848	• Chen et al. (Chen et al., 2021b)	• SARD (Lin et al., 2019)	2881
2849	• MBPP, MathQA-Python (Austin et al., 2021)	• Spider (Yu et al., 2018)	2882
2850	• HumanEval (Chen et al., 2021a)	• JuICe (Agashe et al., 2019)	2883
2851	• APPS (Hendrycks et al., 2021)		

2884 • Nguyen et al. (Nguyen et al., 2019)

2885 • Lin et al. (Lin et al., 2021)

2886 • Zhou et al. (Zhou et al., 2019)

2887 • CoSQL (Yu et al., 2019a)

2888 • SParC (Yu et al., 2019b)

2889 • Malik (Malik et al., 2019)

2890 • LeClair (LeClair et al., 2019)

2891 **2018:**

2892 • CoNaLa (Yin et al., 2018)

2893 • DeepCom (Hu et al., 2018a)

2894 • TL-CodeSum (Hu et al., 2018b)

2895 • code-summarization-public (Wan et al., 2018)

2896 • Russell et al. (Russell et al., 2018)

2897 • VulDeePecker (Li et al., 2018b)

2898 • Lin et al. (Lin et al., 2018)

2899 • StaQC (Yao et al., 2018)

2900 • Advising (Finegan-Dollak et al., 2018)

2901 • ConCode (Iyer et al., 2018)

2902 • NNGen (Liu et al., 2018)

2903 • Gu et al. (Gu et al., 2018)

2904 **2017:**

2905 • QuixBugs (Lin et al., 2017)

2906 • the DeepFix dataset (Gupta et al., 2017)

2907 • Barone et al. (Barone and Sennrich, 2017)

2908 **2016:**

2909 • CODE-NN (Iyer et al., 2016)

2910 • Mou et al. (Mou et al., 2016)

2911 **2015:**

2912 • MANYBUGS, INTROCLASS (Le Goues et al.,
2913 2015)

2914 **2014:**

2915 • Defects4j (Just et al., 2014)

2916 • BigCloneBench (Svajlenko et al., 2014)

E Guideline

Finally, for ease of printing and use, we organized the guideline HOW2BENCH into a clear, color-coded checklist (4 pages in total) that is easy to print, attached at the end of the paper.

2917
2918
2919
2920
2921

No	HOW-TO-BENCH (1/4)	✓
 Phase 0. Benchmark Design		
1	Consider whether the benchmark can fill the gap in related research .	
2	Consider what is the expected scope of the benchmark set (e.g., what natural languages, programming languages, task granularity).	
3	Consider the expected application scenario of this benchmark (e.g., programming assistant, automated tester).	
4	Consider the LLMs' capabilities (e.g., understanding, reasoning, calculation) and domain knowledge (e.g., OOP, memory management, fault localization, process scheduling) that the benchmark hopes to evaluate .	
 Phase 1. Benchmark Construction		
5	Consider whether the data source of the benchmark is traceable .	
6	Consider whether the data source of the benchmark is of high quality (e.g., stars, downloads, last update times, number of forks).	
7	Consider whether the benchmark's data source is representative (e.g., choose an open-source community or code hosting platform that matches the task, capability, and scope under study).	
8	Consider data contamination issues during the benchmark collection (e.g., considering the upload time of the source code or checking whether the data source is included in the training data of LLMs).	
9	If data sampling is needed, consider whether the choice of sample size is scientific (e.g., considering the confidence level/margin of error/sampling proportion, etc.).	
10	If data sampling is needed, consider whether the sampling process is rigorous (e.g., random sampling, stratified sampling, etc.).	
11	Ensure each data point in the benchmark falls into the targeted scope (e.g., checking each data point's evaluated capabilities or domain knowledge).	
12	Consider whether the data in the benchmark can cover the studied capabilities/domain knowledge/application scenarios.	
13	Consider whether there is a standard answer for each sample in the benchmark (such as reference code, etc.).	
14	For code , consider whether the code is compilable/executable .	

No	HOW-TO-BENCH (2/4)	✓
15	Consider the possibility of noise in the data and perform <u>denoise</u> .	
16	Consider the possibility of duplication in the data and <u>deduplicate</u> them.	
17	<u>Clean the sensitive information</u> (such as data desensitization and anonymization) unless the benchmark is deliberately designed so.	
18	<u>Manually review</u> some or all of the data in the benchmark to ensure its quality.	
19	<u>Use LLMs to review</u> some or all of the data in the benchmark to ensure its quality.	
20	Design appropriate <u>output validation methods</u> for the benchmark (e.g., using exact matching or designing test cases).	
21	Design appropriate <u>evaluation metrics</u> for the evaluation set (e.g., precision, accuracy, pass@K, recall).	
22	Consider <u>the adequacy of the evaluation metrics</u> (e.g., is the code coverage high enough).	
23	Consider if there are any <u>other evaluation perspectives</u> (e.g., readability, efficiency, safety, security).	
 Phase 2. Benchmark Evaluation		
24	Consider whether <u>sufficient</u> LLMs are evaluated.	
25	Consider whether <u>representative</u> LLMs (e.g., covering latest/classical LLM families, small/large LLMs, and open-/closed-source LLMs) are evaluated.	
26	Consider whether the prompt is of <u>high quality</u> (e.g., the instruction and intent are clear).	
27	The prompts have been <u>validated by humans or LLMs</u> (e.g., evaluated or discussed by participants or preliminarily tried out on several LLMs).	
28	Try <u>different paraphrases</u> of the prompt.	
29	Try <u>different prompting strategies</u> to observe the impact on the evaluation results (e.g., in-context learning, chain-of-thought).	
30	Pay attention to the <u>hardware environment</u> (such as GPU card, storage size, etc.) of the experiment.	
31	Pay attention to the <u>operating system and software environment</u> (e.g. operating system, version, etc.) used for the experiment.	

No	HOW-TO-BENCH (3/4)	✓
32	Pay attention to the off-the-shelf <u>platforms, frameworks, or libraries</u> for LLM evaluation (e.g., fast chat, vllm, huggingface) that are used.	
33	<u>Repeat</u> the experiment multiple times to reduce the impact of <u>randomness</u> on the evaluation.	
34	Consider various <u>randomization strategies</u> (e.g., trying various temperature parameters) to reduce the impact of parameter configuration on the evaluation.	
35	<u>Record</u> the experimental process in detail (e.g., parameter settings, running time, input/output pairs, etc.).	
 Phase 3. Benchmark Analysis		
36	Observe the <u>difficulty</u> of the benchmark, checking if the benchmark is too hard or too easy for LLMs (i.e., most LLMs score too high/low).	
37	Consider whether the benchmark can <u>distinguish</u> the pros and cons of different LLMs.	
38	If the experiment is <u>repeated several times</u>, consider the <u>stability</u> of the benchmark (i.e., whether the experimental results vary too much in the repeated experiments).	
39	Analyze the <u>correlation</u> between the data and their score. For example, if there is a correlation between the data (such as similar difficulty and knowledge required), then the scores should also be correlated.	
40	Compare the performance of LLMs on this benchmark with their performance on other related benchmarks.	
41	Consider <u>presenting</u> the experiment results <u>in an appropriate way</u> (e.g., table, line graph, pie chart, etc.).	
42	Consider <u>presenting</u> the experiment results <u>clearly</u> (e.g., distinguishable colors/labels/shapes, etc.).	
43	<u>Explain</u> the experiment results.	
44	Observe <u>correlations via multiple perspectives</u> from the experimental results (e.g., performance is correlated with model size or amount of context).	
45	The analysis of the evaluation results will be <u>inspiring</u> (e.g., shed light on future direction, make actionable advice, etc.).	

No	HOW-TO-BENCH (4/4)	✓
 Phase 4. Benchmark Release		
46	Set the appropriate <u>license</u> for the benchmark.	
47	Review the released benchmark or other artifacts to ensure they do NOT contain sensitive information (e.g., API keys, usernames, passwords, etc.).	
48	review the released benchmark or other artifacts to ensure they do NOT contain toxicity information (e.g., abusive comments/identifiers).	
49	Make sure the benchmark is <u>open-accessible</u> .	
50	Make sure the <u>test cases</u> or <u>reference data</u> are open and accessible.	
51	Provide <u>prompts</u> used in the experiment to ensure the experiments are reproducible.	
52	Disclose the experimental environment (e.g., hardware, operating system, software version, framework platform) to ensure the reproducibility of the experiment.	
53	Make the <u>detailed</u> experimental results <u>public</u> for verification.	
54	Ensure the <u>quality of the user manual</u> such as README (e.g., it contains necessary benchmark introduction, executable scripts, etc.).	
55	Provide <u>convenient evaluation interfaces</u> for the released benchmark (e.g., providing a command line interface, docker, etc.).	