
BetterBench: Assessing AI Benchmarks, Uncovering Issues, and Establishing Best Practices

Anka Reuel *
Stanford University

Amelia Hardy*
Stanford University

Chandler Smith
Northeastern University

Max Lamparth
Stanford University

Malcolm Hardy
Stanford University

Mykel J. Kochenderfer
Stanford University

Abstract

1 AI models are increasingly prevalent in high-stakes environments, necessitating
2 thorough assessment of their capabilities and risks. Benchmarks are popular for
3 measuring these attributes and for comparing model performance, tracking progress,
4 and identifying weaknesses in foundation and non-foundation models. They can
5 inform model selection for downstream tasks and influence policy initiatives.
6 However, not all benchmarks are the same: their quality depends on their design
7 and usability. In this paper, we develop an assessment framework considering 46
8 best practices across an AI benchmark’s lifecycle and evaluate 24 AI benchmarks
9 against it. We find that there exist large quality differences and that commonly used
10 benchmarks suffer from significant issues. We further find that most benchmarks
11 do not report statistical significance of their results nor allow for their results to be
12 easily replicated. To support benchmark developers in aligning with best practices,
13 we provide a checklist for minimum quality assurance based on our assessment. We
14 also develop a living repository of benchmark assessments to support benchmark
15 comparability, accessible at betterbench.stanford.edu.

16 1 Introduction

17 AI systems are rapidly advancing and proliferating [58]. The increasing integration of AI, and in
18 particular foundation models (FMs) [14], into decision-making systems has significantly amplified
19 its impact and has showcased both benefits [9, 39, 57, 66] and risks [2, 75, 44, 86, 45, 30, 70]. Given
20 the importance of correctly assessing a model’s capabilities and potential harms, AI evaluation is
21 an essential discipline [15]. Current evaluation approaches include both internally (e.g., private
22 testing on proprietary data) and externally developed techniques (e.g., scoring on public benchmarks)
23 [74, 27, 73, 48, 32].

24 Following the work of [67], we define a benchmark “as a particular combination of a dataset or sets
25 of datasets [...], and a metric, conceptualized as representing one or more specific tasks or sets of
26 abilities, picked up by a community of researchers as a shared framework for the comparison of
27 methods” [67]. Using benchmarks to facilitate comparison, measure performance, track progress, and
28 identify weaknesses has become a standard practice. For example, benchmarks are widely used by
29 model developers to report performance and compare models upon release [3, 8], and as part of policy
30 initiatives to support third-party model evaluations, such as as part of the UK AI Safety Institute’s

() denotes equal contribution. Corresponding authors: anka.reuel@stanford.edu, ahardy@stanford.edu

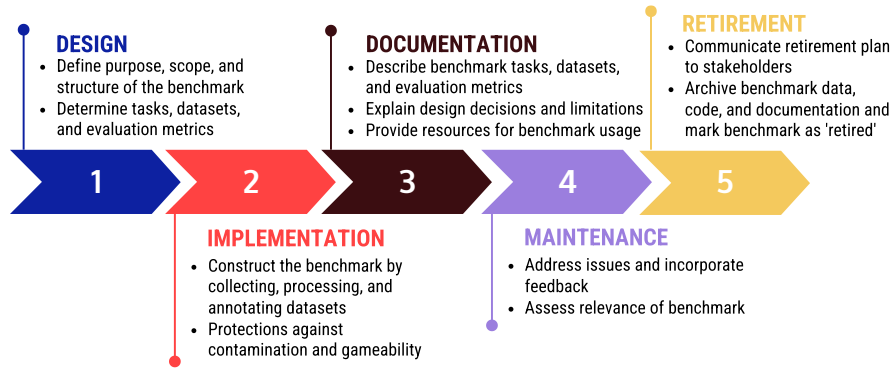


Figure 1: Five stages of the benchmark lifecycle. A detailed description can be found in App. C.

Inspect framework for evaluating large language models (LLMs) [81] or Article 51 of the EU AI Act [1]. However, the fidelity of this approach depends entirely on the benchmarks' quality, where we define a *high-quality* benchmark as one that is interpretable, clear about its intended purpose and scope, and that is usable. To date, no structured assessment for the quality of AI benchmarks, including both FM and non-FM benchmarks, has been published, and no comparative analysis has been conducted to understand quality differences between widely used AI benchmarks. To address these gaps, our paper:

- Presents a novel AI benchmark assessment framework evaluating the quality of AI benchmarks based on 46 criteria derived from expert interviews and domain literature
- Scores 16 foundation model (FM) and 8 non-FM benchmarks (full list in App. D), finding quality differences across both categories
- Provides insights into prevalent issues in current AI benchmarking practices based on our assessment
- Creates a checklist for minimum quality assurance to support benchmark developers in aligning with best practices
- Makes available a living repository² of benchmark assessments for users to analyze benchmarks' quality and appropriateness for their usage contexts.

We structure the paper as follows: Sec. 2 explores benchmarking in AI and other fields. Sec. 3 describes our assessment development, which combined literature and expert interviews, and details our benchmark scoring procedure. Sec. 4 presents our framework's criteria, focusing on aspects under developers' control to promote better benchmarks. Sec. 5 lists additional context-dependent design considerations. Sec. 6 reports findings from applying our framework to 24 benchmarks. Finally, Sec. 7 and Sec. 8 explore implications for future evaluations and discuss our work's scope and limitations. We further outline open challenges with AI benchmarking in App. A, involved stakeholders in App. B, and the AI benchmark lifecycle in App. C.

2 Related Work

2.1 AI Benchmarking Practices and Challenges

Our literature review of AI benchmarking practices identifies two primary concerns: what a benchmark measures and how this measurement is used. Regarding what a benchmark measures, [59] find that current benchmarks for LLMs are insufficient for assessing these models capabilities. A frequent concern in this context is the validity of evaluations [54, 76, 67]. Similarly, [62] finds

²<https://betterbench.stanford.edu>

62 that the rapid advancement of AI models threatens benchmarks' utility, as a large fraction of these
63 evaluations are near saturation. [83] and [49] both address the narrow scope of existing benchmarks,
64 with [49] advocating for approaches intended to reduce the socio-technical gap that exists between
65 the capabilities that benchmarks are able to measure and the ability of models to meet user needs
66 in downstream applications. With respect to how evaluations are used, [67] critiques the tendency
67 of AI practitioners to overgeneralize benchmark results, highlighting how these scores present an
68 inherently reductive view of model performance.

69 In addition, the community has also recognized the importance of data curation and documentation
70 in the context of evaluations. [65] put forth the idea of data cards as standardized documentation
71 framework for datasets and [12] develop a framework and checklist for best practices in data curation.
72 Finally, the FAIR principles [87] outline best practices for digital data access, based on the principles
73 of *Findability*, *Accessibility*, *Interoperability*, and *Reuse*. While these efforts support the adoption
74 of best practices in the context of data, they are insufficient for assessing AI benchmarks, which
75 extend data with infrastructure and evaluation methods, requiring additional guidelines to support the
76 development of high-quality benchmarks and the decision-making of benchmark users.

77 Hence, our work builds on and expands these guidelines, with the aim of advancing the analysis of
78 AI benchmarking by presenting a first-of-its-kind framework for the assessment of both foundation
79 model and non-foundation model benchmarks. Unlike prior studies, such as [59] and [49], which
80 focus on identifying limitations in limited contexts and scopes, our approach offers practical tools,
81 empowering developers to address shortcomings and directly enhance benchmark quality: Our
82 assessment spans a wider range of criteria across the benchmark lifecycle, from design (e.g., have
83 domain experts been involved in the development?) to implementation (e.g., is the evaluation script
84 available?), documentation (e.g., is the applicable license specified?), and maintenance (e.g., is a
85 feedback channel available for users?). We give an overview of all our criteria in Sec. 4 and explain,
86 justify, and provide scoring details for each criterion in App. K. We further provide a checklist of best
87 practices derived from our analysis (App. J), offering guidance for improving AI benchmarks, rather
88 than merely highlighting issues.

89 2.2 Benchmarking Best Practices in Other Fields

90 Our work is informed by benchmarking practices from fields beyond AI, ranging from transistor
91 hardware [18] to environmental quality [16] to bioinformatics [7], and we identify common themes
92 regarding what constitutes an effective benchmark. Where applicable, we incorporate these best
93 practices into our assessment (Sec. 4):

94 **Designing for downstream utility.** Many of the papers reviewed discuss the importance of a
95 benchmark's tasks being designed with real world applications in mind. [16] considers the best
96 benchmarks to be situation-specific, [24] defines an ideal test set as one which reflects real world data,
97 [7] proposes that benchmarks should be adapted to their intended applications, and [25] suggests
98 that benchmarks be designed to fit the diversity of downstream use cases. [77] emphasizes the
99 importance of guaranteeing that tested methods only use information available in a practical setting
100 and recommends checking that a benchmark simulates the envisioned usage.

101 **Ensuring validity.** A frequent concern with benchmarking is the validity of evaluations [54, 76, 67].
102 In educational testing, [60] outline a framework to ensure validity by providing guidelines for effective
103 evidence collection. [22] outline what and how evidence can be collected and how it should be
104 interpreted for tests "of attributes for which there is no adequate criterion" [22]. Measures that are
105 used in other fields further include choosing a large test set to promote the statistical significance of
106 results [77] and updating a benchmark over time to prevent developers from overfitting it [7]. [7] also
107 notes that the methods or approaches being evaluated should not be used to create the gold standard
108 dataset.

109 **Prioritizing score interpretability.** [7] highlights that benchmarks are particularly important when
110 a wide variety of tools are available and it is difficult for non-specialists to distinguish between
111 them. Interpretability is important in not only selecting tools, but also deciding between benchmarks

112 themselves. Effective benchmarks must provide transparent information regarding the procedural
113 details of their experiments [18] and goals of the evaluation [10]. They should clearly describe the
114 benchmark’s purpose and scope, as these are fundamental to its design and implementation [85].
115 Regarding scope, [16] states that for environmental quality applications, benchmarks should never be
116 the basis of final decisions. With this in mind, they identify misleading benchmarks as the worst-case
117 scenario. Furthermore, they state that a benchmark should not present its results as absolutes, instead
118 ensuring that its evaluations are understandable inputs for decision makers [16].

119 **Guaranteeing accessibility.** A good benchmark is easy to obtain and use [7, 77, 25, 10]. If a
120 benchmark is run computationally, then its data and scripts must be available for results to be
121 reproducible [77, 25, 10].

122 3 Methodology

123 Our benchmark assessment consists of 46 criteria based on our literature review and interviews
124 with five primary groups of stakeholders. These groups, who also present the user personas of our
125 assessment, are described in detail in App. B. Through our interview process, we defined a five-stage
126 benchmark lifecycle and identified objectives along it. In this section, we discuss our methodology
127 for identifying stakeholders, developing criteria, and assessing benchmarks. A detailed flow diagram
128 of our methodology can be found in App. H.

129 **Step 1: Mapping the space.** Initially, we surveyed the existing benchmark landscape (Sec. 2).
130 Based on this review, we identified five stakeholder groups who present the user personas of our
131 assessment (App. B). To understand their objectives with respect to benchmarking, we conducted
132 unstructured interviews with representatives of all stakeholder groups, including 20+ policymakers,
133 model developers, benchmark developers, model users, and AI researchers. During this process, we
134 developed a five-stage model of the benchmark lifecycle (Fig. 5 and App. C) and mapped both the
135 benchmarking objectives of the stakeholders and their communicated use cases for a benchmark
136 assessment (App. B).

137 **Step 2: Translation to criteria.** Based on Step 1, we identified tasks and objectives for each stage
138 of the AI benchmark lifecycle and translated them into concrete criteria. We categorized these
139 as: (a) criteria controlled by the benchmark developer where the authors and interviewees reached
140 a normative consensus, (b) criteria controlled by the benchmark developer but context-dependent,
141 difficult for an external party to assess, or both and (c) aspects either outside the benchmark developer’s
142 control or requiring further research. The assessment in Sec. 4 is limited to category (a) criteria. We
143 cover considerations in (b) in Sec. 5, and those in (c) in App. A.

144 **Step 3: Validating the assessment.** Initially, three authors independently scored the same benchmark
145 to calibrate the assessment and identify potential misinterpretations of the criteria. We adapted and
146 clarified scoring guidelines (App. K) to address differing interpretations and uncertainties. To validate
147 our assessment, we shared it with members of all stakeholder groups and revised it based on their
148 feedback. Finally, we verified that our assessment, which in itself can be considered a benchmark,
149 met all of our defined criteria, where applicable (App. J.2).

150 **Step 4: Structuring the assessment.** We evaluated 16 FM and 8 non-FM benchmarks. We prioritized
151 commonly used benchmarks, such as those that were recently reported by model developers [8, 3]
152 and aim to expand the number of assessed benchmarks continuously on our website *betterbench.stan-*
153 *ford.edu*. Since our assessment considers varying information sources (official websites, papers,
154 GitHub repositories published by the benchmark developers³) that do not follow a standard structure,
155 we manually evaluated all benchmarks. At least two authors independently reviewed each benchmark.
156 They subsequently had to reach a consensus on the final score and a third reviewer could be called to
157 make the final decision if a consensus could not be reached (this case did not occur).

³We do not consider third-party information that was not released by the benchmark developers themselves.

Step 5: Scoring. We scored benchmarks on a discrete 0/5/10/15-point scale for each criterion: 15 for fully meeting, 10 for partially meeting, 5 for mentioning without fulfilling, and 0 for neither referencing nor satisfying the criterion. Average scores were calculated for each benchmark lifecycle stage (design, implementation, documentation, and maintenance). An aggregate usability score, representing the weighted average of the implementation, documentation, and maintenance scores, was also introduced (see App. G for scoring details). We consider a mean score of 10 or higher to indicate a reasonably good benchmark for each aggregated scoring category, as it signifies that, on average, the benchmark at least partially fulfills all assessment criteria within the respective category.

Step 6: Platform for continuous updates. Finally, we develop a supplementary website⁴ to continuously publish assessment results using the scoring methodology in App. G, given the rapid development of new AI benchmarks. The website includes a community feedback channel for submitting new AI benchmarks and correcting previously posted scores if benchmarks are updated or stakeholders disagree with our evaluation. This provides benchmark users with an accessible, up-to-date database of existing benchmarks and their quality, enabling quick analysis of the most suitable benchmark for their application context.

4 Assessment Criteria

We separate our assessment criteria according to the phase of the benchmark lifecycle during which they would be fulfilled. Although the retirement stage is within the developer’s control, we do not include specific criteria for this phase within the current framework, because we cannot assess the retirement of active benchmarks. App. K contains full explanations, justifications, and scoring guidelines for each of the 46 criteria.

4.1 Benchmark Design

Design Criteria	
1. Tested capability, characteristic, or concept is defined	9. How benchmark score should or shouldn't be interpreted or used is described
2. How tested capability or concept translates to benchmark task is described	10. How knowing about the tested concept is helpful in the real world is described
3. Domain experts are involved	11. Informed performance metric choice
4. Domain literature is integrated	12. Metric floors and ceilings are included
5. Use cases or user personas are described	13. Human performance level is included
6. Differences to related benchmarks are explained	14. Random performance level is included
7. Input sensitivity is addressed	
8. Has validated automatic evaluation	

Figure 2: Overview of assessment criteria for the benchmark design stage.

Benchmarks should clearly describe their goals and scope [85, 10, 54]. This includes defining the tested capability or characteristic, describing how the tested capability translates to the benchmark task, and stating how knowing about the tested concept is helpful in real-world applications [54]. These design choices should be informed by considering use cases and user personas for the benchmark, involving domain experts, and integrating domain literature [82]. Clearly stating how the benchmark is different from related existing AI benchmarks is necessary to help benchmark users decide the applicability of a benchmark to their use case. A benchmark’s measurements must be interpretable [16], which requires an informed choice of performance metric(s) and a description of how the benchmark score should or shouldn’t be interpreted [48]. Including floors, ceilings, human performance levels, and random performance levels for the chosen metric(s) further assists users in understanding a model’s score [34]. If addressing input sensitivity and providing a validated automatic evaluation are possible, these measures enhance a benchmark’s robustness and accessibility [34].

⁴betterbench.stanford.edu. Our assessment and results are released under a CC BY 4.0 license.

193 **4.2 Benchmark Implementation**

Implementation Criteria	
1. Evaluation code is available	7. Script to replicate results is explicitly included
2. Evaluation data or generation mechanism is accessible	8. Statistical significance or uncertainty quantification of benchmark results is reported
3. Evaluation of models via API is supported	9. Need for warnings for sensitive/harmful content is assessed
4. Evaluation of local models is supported	10. Build status is implemented
5. Globally unique identifier or encryption of evaluation instances is added	11. Release requirements are specified
6. Task to identify if model has been trained on benchmark data is included	

Figure 3: Overview of assessment criteria for the benchmark implementation stage.

194 Criteria in the implementation stage focus on the availability of necessary code and infrastructure
195 and the inclusion of key engineering features. To ensure reproducibility and scrutiny [77, 25, 10],
196 a benchmark should provide working evaluation code, and make its evaluation data, prompts, or
197 dynamic test environment accessible. A script should be available to replicate initial published
198 results. In domains where models are often accessed via API, such as NLP, an ideal benchmark
199 supports the evaluation of both API-based and local models. A benchmark can minimize the risks of
200 contamination and gamification by including a globally unique identifier or encrypting evaluation
201 instances. This is especially important for testing models that rely on web-scraped training data.
202 Including a *training_on_test_set* task allows determining whether a model’s training data included
203 benchmark examples [74]. As an additional measure, specifying clear release requirements informs
204 users how to preserve the integrity of test results [6].

205 **4.3 Benchmark Documentation**

Documentation Criteria	
1. Requirements file available or equivalent is available	11. Globally unique, persistent identifier for a dataset and its metadata is provided
2. Quick-start guide or demo is available	12. Standardized metadata is included
3. In-line code comments are used	13. Data sources and data collection process are explained
4. Code documentation is available	14. Data preprocessing steps are described (if applicable)
5. Accompanying paper is accepted at peer-reviewed venue	15. Data annotation process is described (if applicable)
6. Benchmark design process is documented	16. Evaluation metric is documented
7. Test tasks & rationale are documented	17. Applicable license is specified
8. Assumptions of normative properties are documented	18. Data representativeness is explained (if applicable)
9. Limitations are documented	19. Data is documented using a standardized format.
10. Test environment design or prompt design process is documented	

Figure 4: Overview of assessment criteria for the benchmark documentation stage.

206 Providing comprehensive and accessible documentation is crucial for the practicability and interpreta-
207 tion of benchmarks [18]. Key information about a benchmark should be readily available and include
208 documentation of benchmark construction processes [54], data collection [87] or test environment
209 design, and its test tasks and their rationale [54]. Clearly documenting evaluation metric(s) and
210 reporting the statistical significance of results is necessary so that users can understand a benchmark’s
211 actual signal [4]. To provide context and prevent misinterpretation, developers should document
212 normative assumptions about benchmark properties and discuss the limitations of their benchmark.
213 A benchmark’s codebase should contain a requirements file, a quick-start guide or demo code, a
214 description of code file structure and contents, and in-line comments within all relevant files. Having
215 a benchmark’s paper accepted at a peer-reviewed venue signals external scrutiny and adherence to
216 certain standards. Lastly, developers should specify the applicable license to provide legal clarity and
217 enable, e.g., commercial use.

Maintenance Criteria	
1. Code usability was checked within the last year 2. Maintained feedback channel for users is available	3. Contact person is listed

Figure 5: Overview of assessment criteria for the benchmark maintenance stage.

219 An optimally designed, implemented, and documented benchmark will cease to be useful if it is not
220 maintained. Developers should regularly check code usability and maintain a feedback channel for
221 users to report issues or suggest improvements. Providing contact details of a person responsible for
222 the benchmark facilitates communication and support. Alternatively, if a benchmark is not maintained
223 anymore, authors should include a corresponding statement indicating that the benchmark was retired
224 in any official benchmark artefacts.

225 5 Other Design Considerations

226 This section presents design considerations for benchmark developers that were excluded from our
227 assessment because their appropriateness is context-dependent, they are not easily verifiable, or both.
228 Our aim with this list is to promote conscious design decisions regarding these considerations.

229 **General vs. specific benchmarks.** Benchmark developers must decide whether to prioritize general
230 or abstract knowledge and skills or specific contexts and domains. Broad concept benchmarks may
231 contribute to understanding foundational characteristics of models, but often face challenges in
232 real-world applicability and reliable testing (see App. A).

233 **Detecting small improvements.** Benchmarks should be designed so that a 1% improvement can be
234 reliably detected [34]. As [34] states, “the more difficult it is to detect small amounts of progress,
235 the more difficult it becomes to make iterative progress on a benchmark.” Practically, this is likely
236 dependent on evaluation data size and task diversity.

237 **Multi-modal assessment.** As multi-modal models become increasingly common, benchmark de-
238 velopers may want to consider designing tasks to assess the capabilities they want to test across
239 modalities. Additional design considerations for multi-modal assessments include the increased
240 complexity of mapping a tested concept to different modalities and the different output formats of the
241 tested models [91].

242 **Versioning.** Minor updates (e.g., removing faulty prompts) should be clearly indicated via *task*
243 *versioning* [13]. Major updates require releasing new *benchmark versions*, as exemplified by the
244 AgentBench v0.1 and v0.2 releases [52].

245 **Dynamic vs. static benchmarks.** Dynamic benchmarks may better address quick saturation (App. A)
246 and contamination (App. A) issues but reduce result comparability and are easier to implement for
247 some tasks (e.g., adding numbers) than others. Static benchmarks, on the other hand, tend to suffer
248 from the issues outlined above.

249 **Gameability.** An ideal benchmark is resilient to attempts to boost task performance without im-
250 proving the fundamental capability being tested [7]. Existing benchmarks have been shown to be
251 vulnerable to manipulation [6]. Specific guidelines have been proposed to prevent cheating and
252 ensure evaluations reflect genuine model performance [94].

253 **Positionality statement.** Positionality statements⁵ are a reflective account common in social sciences
254 research. In them, researchers acknowledge how their background, experiences, and biases may have
255 influenced their work. If developers believe such factors significantly impacted their benchmark’s
256 construction, they may provide a positionality statement for increased context and transparency.

⁵Such statements were not included in the assessment to avoid pressuring benchmark developers to disclose potentially sensitive personal information, even if such information influenced the benchmark design process.

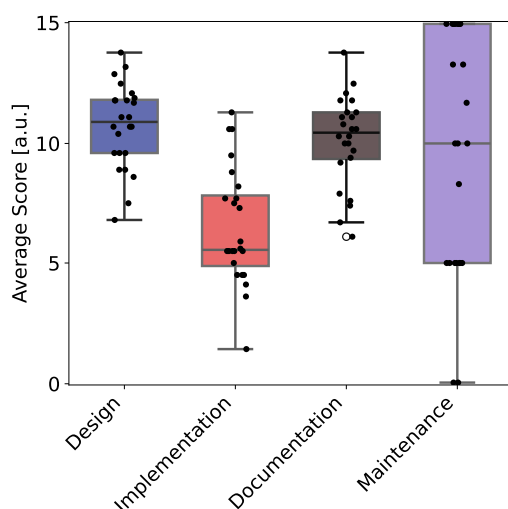


Figure 6: Average and individual scores of all assessed benchmarks per lifecycle stage.

Stage	FM	Non-FM	All
Design	10.6	11.1	10.7
Implementation	5.5	7.4	6.1
Documentation	10.3	9.9	10.1
Maintenance	9.1	10.8	9.7

Table 1: Benchmark lifecycle scores averaged over the 24 assessed benchmarks separated for FM, non-FM, and All benchmarks combined.

	FM	Non-FM	All
Pearson ρ	0.721	0.318	0.655
p-value p	0.001	0.487	0.001

Table 2: Pearson correlation coefficient for FM, Non-FM, and All benchmarks between the design and usability (weighted average of implementation, documentation, and maintenance stages) score as in Fig. 7.

6 Quantitative Results

In this section, we present our assessment results.⁶ Tab. 1 showcases the average scores per benchmark lifecycle stage, showing that for both FM and non-FM benchmarks, the implementation stage tends to be the weakest area, followed by maintenance. All criteria averages are reported in App. F. Some criteria have not been fulfilled by almost any benchmark (e.g., *Standardized metadata is included*). Notably, both benchmark types are particularly weak for criteria supporting the reproducibility and interpretation of results: benchmarks get an average score of 3.75 on *Including a script to replicate results* and an average score of 5.62 on *Reporting statistical significance*.

While individual benchmark or criteria scores are deterministic, we can analyze statistical fluctuations across categories and benchmarks. Fig. 7 compares the design and usability scores of FM and non-FM benchmarks. The overall average design score across all benchmarks is 10.7, and the weighted average usability score is 8.7. The difference in mean design and usability scores between FM and non-FM benchmarks is not statistically significant (95% confidence level), see Fig. 8 in App. E. Furthermore, we find statistically significant correlations between the design and usability scores for FM benchmarks alone and all benchmarks combined at the 95% confidence level (Tab. 2). This suggests that, in both cases, benchmarks with poorer design tend to also be less usable, and vice versa.

7 Discussion

Not all benchmarks are of the same quality. Model developers frequently report performance on benchmarks that vary significantly in quality. For instance, the widely-used MMLU benchmark scored the lowest in our assessment (weighted average: 5.5), while GPQA scored significantly higher (weighted average: 11.0). However, recent communications introducing models like GPT-4 [3], Claude-3 [8], and Gemini [80] report results on both benchmarks without explicitly acknowledging their limitations or quality differences. This practice may be driven by the assumed expectation that reviewers want to see a wide range of metrics and the belief that readers should determine the most relevant metrics for their needs. The lack of clear guidance on AI benchmark quality and limitations may lead to incorrect conclusions about a model’s performance, even if developers do not intend to

⁶Per-criterion scores for all benchmarks are released on our website betterbench.stanford.edu. Code to replicate results will be available on GitHub upon publication.

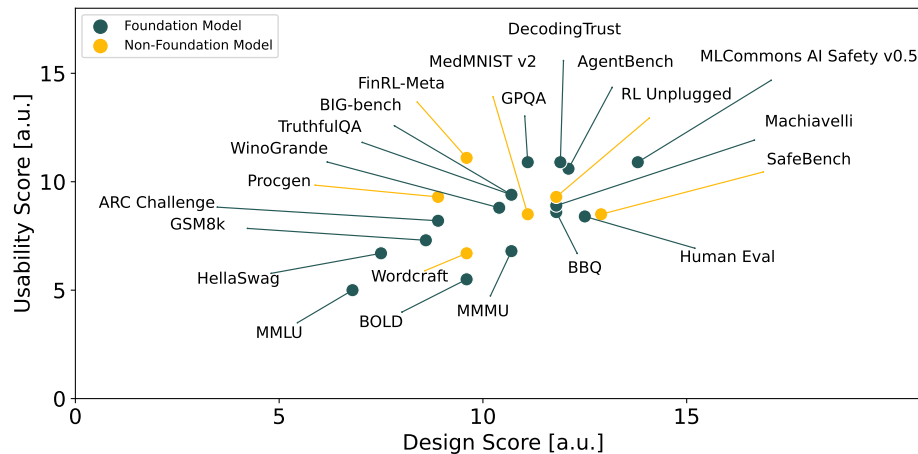


Figure 7: Design and usability score for all 24 assessed benchmarks. The usability score is the weighted average of the implementation, documentation, and maintenance scores. Benchmarks were split into foundation model and non-foundation model benchmarks, depending on the model group they're targeting.

mislead users. The UK AI Safety Institute's *Inspect* framework [81] similarly includes both MMLU [33] and GPQA [68], potentially resulting in misleading evaluations. This is problematic because governments increasingly rely on evaluations for AI regulations and may use frameworks like *Inspect* [69] or individual benchmarks [1].

Most benchmarks fail to distinguish signal and noise. Benchmark developers should not only report a single result for a model but also re-run their evaluation [13] with, e.g., different random seeds or sampling temperatures, and report the mean and variance for these intra-model evaluations. As benchmarks are primarily used to compare models, users must know the intra-model variance of a benchmark to determine whether observed inter-model variances are genuine performance differences or arise from noisy results. If intra-model variance bounds are tight and inter-model variance bounds are wide, benchmark users can conclude that there are genuine performance differences between models. However, if both intra- and inter-variance bounds are wide, statistical analysis is required to discern noise and actual signal. Yet, 14 out of 24 benchmarks did not perform multiple evaluations of the same model or report statistical significance or uncertainty of results.

Insufficient implementation limits reproducibility and scrutiny of benchmarks. Our analysis reveals that scores for implementation stage criteria are the lowest across all assessed benchmarks. Notably, 17 out of 24 benchmarks do not provide easy-to-run scripts to replicate the results reported in the initial paper, and 4 out of 24 only provide scripts to replicate part of the results. This lack of accessibility hinders reproducibility and limits users' ability to scrutinize the benchmarking process. In a field where reproducibility is a significant concern [43], providing materials to reproduce results is crucial for validating benchmark findings.

Small changes can lead to significant improvements in overall benchmark practices. Many of the criteria we have identified for improving AI benchmarks are relatively easy to implement, even for existing benchmarks. For example, adding code documentation and a point of contact are not time consuming to add, yet can significantly enhance usability, accountability, and ease of use.

Necessity for higher benchmark development standards. As evidenced by the strong discrepancies in AI benchmark quality we found (Sec. 6 and App. F), there is a need to introduce additional checks for benchmarking practices to ensure a minimum quality standard for AI benchmarks. We assume that benchmark developers do not intentionally construct insufficient benchmarks, but rather do so due to limited knowledge of what constitutes a good benchmark. By providing a checklist of best practices (App. J.1), we aim to make it easy for benchmark developers to adopt these recommendations and

improve the quality of their benchmarks. In addition, some of the criteria we have identified in our expert interviews and from reviewing evaluation practices in other fields, such as including a build status in GitHub repositories that assesses whether the last commit successfully passed defined unit tests [28], were relatively unknown and only implemented by 3 out of 24 benchmarks. Other criteria, like using globally unique identifiers or encrypting evaluation instances to avoid data contamination, have been pioneered by only a few of the assessed benchmarks [68, 74] but have not yet gained widespread adoption. By incorporating these criteria into our assessment, we aim to encourage benchmark developers to adopt these best practices in the field of AI benchmarking.

8 Limitations

Our assessment assigns equal weight to all criteria, despite their varying levels of effort required for fulfillment and differing contributions to overall benchmark quality. The scoring system differentiates only four score categories to enable relatively objective evaluation through clear-cut criteria (App. K and App. G), but may miss nuances within each category. For example, a benchmark barely fulfilling a criterion and one almost entirely fulfilling it would receive the same 10-point score. Given the equal weighting and scoring, benchmark developers could potentially “game” the assessment by focusing on easily fulfilled criteria. However, we believe that even if a developer only implements easy-to-implement criteria, the resulting benchmark will still be of higher quality than one not meeting any criteria, thus fulfilling our work’s goal. Furthermore, assessing the construct validity of a benchmark and determining whether its approach to assessing a concept is truly effective would presumably require in-depth analysis by domain experts in the respective fields, which is beyond the scope of this assessment. Instead, we aim to provide benchmark developers with a blueprint for minimum quality assurances. Finally, our framework is intended for public benchmarks and future work is needed to extend it to private ones.

9 Impact Statement

By releasing the first systematic assessment framework for AI benchmarks, we aim to encourage benchmark developers to construct higher-quality benchmarks and to contribute to community efforts to make AI evaluations more practicable and transparent. Higher-quality benchmarks resulting from the adoption of our framework and checklist can lead to better-informed model selection for downstream tasks, potentially reducing risks and improving outcomes in high-stakes applications. Our living repository of benchmark assessments promotes transparency and comparability, allowing benchmark users to make informed decisions when choosing benchmarks. However, there is a potential risk of misinterpretation of our results; our assessment only provides minimum quality assurances and is not sufficient to assess the suitability of a benchmark for a concrete use case. The outputs of our evaluation do not contain sensitive or harmful content, but users may encounter such content during a benchmark assessment depending on the benchmark’s data. While we do not anticipate direct safety risks from releasing our framework, we acknowledge that strict adherence to some of our proposed criteria, such as the involvement of domain experts, may unequally impact researchers based on their access to resources and connections, potentially hindering the development of benchmarks from a broader range of research institutions and underrepresented communities, which could limit diversity in benchmark creation.

References

- [1] Art. 51 Classification of General-Purpose AI Models as General-Purpose AI Models with Systemic Risk - EU AI Act — euaiact.com. <https://www.euaiact.com/article/51>. [Accessed 31-07-2024].
- [2] Abubakar Abid, Maheen Farooqi, and James Zou. Persistent anti-muslim bias in large language models. In *AAAI/ACM Conference on AI, Ethics, and Society*, pages 298–306, 2021.

- [3] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. GPT-4 Technical Report. *arXiv preprint arXiv:2303.08774*, 2023.
- [4] Rishabh Agarwal, Max Schwarzer, Pablo Samuel Castro, Aaron C Courville, and Marc Bellemare. Deep reinforcement learning at the edge of the statistical precipice. *Advances in neural information processing systems*, 34:29304–29320, 2021.
- [5] Parveen Azam Ali and Roger Watson. Peer review and the publication process. *Nursing open*, 3(4):193–202, 2016.
- [6] Norah Alzahrani, Hisham Abdullah Alyahya, et al. When Benchmarks are Targets: Revealing the Sensitivity of Large Language Model Leaderboards. *ArXiv*, abs/2402.01781, 2024.
- [7] Mohamed Radhouene Aniba, Olivier Poch, and Julie D. Thompson. Issues in bioinformatics benchmarking: the case study of multiple sequence alignment. *Nucleic Acids Research*, 38(21):7353–7363, 07 2010.
- [8] Anthropic. Introducing the next generation of Claude. <https://www.anthropic.com/news/claude-3-family>, 2024. Accessed on June 2, 2024.
- [9] David Baidoo-Anu and Leticia Owusu Ansah. Education in the era of generative artificial intelligence (AI): Understanding the potential benefits of ChatGPT in promoting teaching and learning. *Journal of AI*, 7(1):52–62, 2023.
- [10] Thomas Bartz-Beielstein, Carola Doerr, Jakob Bossek, Sowmya Chandrasekaran, Tome Eftimov, Andreas Fischbach, Pascal Kerschke, Manuel López-Ibáñez, Katherine Mary Malan, Jason H. Moore, Boris Naujoks, Patryk Orzechowski, Vanessa Volz, Markus Wagner, and T. Weise. Benchmarking in Optimization: Best Practice and Open Issues. *ArXiv*, abs/2007.03488, 2020.
- [11] M. G. Bellemare, Y. Naddaf, J. Veness, and M. Bowling. The Arcade Learning Environment: An Evaluation Platform for General Agents. *Journal of Artificial Intelligence Research*, 47:253–279, jun 2013.
- [12] Eshta Bhardwaj, Harshit Gujral, Siyi Wu, Ciara Zogheib, Tegan Maharaj, and Christoph Becker. Machine learning data practices through a data curation lens: An evaluation framework. In *The 2024 ACM Conference on Fairness, Accountability, and Transparency*, pages 1055–1067, 2024.
- [13] Stella Biderman, Hailey Schoelkopf, Lintang Sutawika, Leo Gao, Jonathan Tow, Baber Abbasi, Alham Fikri Aji, Pawan Sasanka Ammanamanchi, Sidney Black, Jordan Clive, Anthony DiPofi, Julen Etxaniz, Benjamin Fattori, Jessica Zosa Forde, Charles Foster, Jeffrey Hsu, Mimansa Jaiswal, Wilson Y. Lee, Haonan Li, Charles Lovering, Niklas Muennighoff, Ellie Pavlick, Jason Phang, Aviya Skowron, Samson Tan, Xiangru Tang, Kevin A. Wang, Genta Indra Winata, François Yvon, and Andy Zou. Lessons from the Trenches on Reproducible Evaluation of Language Models. *arXiv preprint arXiv:2405.14782*, 2024.
- [14] Rishi Bommasani, Drew A Hudson, Ehsan Adeli, Russ Altman, Simran Arora, Sydney von Arx, Michael S Bernstein, Jeannette Bohg, Antoine Bosselut, Emma Brunskill, et al. On the opportunities and risks of foundation models. *arXiv preprint arXiv:2108.07258*, 2021.
- [15] Yupeng Chang, Xu Wang, Jindong Wang, Yuan Wu, Linyi Yang, Kaijie Zhu, Hao Chen, Xiaoyuan Yi, Cunxiang Wang, Yidong Wang, Wei Ye, Yue Zhang, Yi Chang, Philip S. Yu, Qiang Yang, and Xing Xie. A Survey on Evaluation of Large Language Models. *ACM Trans. Intell. Syst. Technol.*, 15(3), mar 2024.
- [16] Peter M Chapman. Environmental Quality Benchmarks — The Good, The Bad, and The Ugly. *Environmental Science and Pollution Research*, 25(4):3043–3046, 2018.

- [17] 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 Hebgen Guss, 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. Evaluating Large Language Models Trained on Code. *arXiv preprint arXiv:2107.03374*, 2021.
- [18] Zhihui Cheng, Chin-Sheng Pang, Peiqi Wang, Son T Le, Yanqing Wu, Davood Shahrjerdi, Iuliana Radu, Max C Lemme, Lian-Mao Peng, Xiangfeng Duan, et al. How to report and benchmark emerging field-effect transistors. *Nature Electronics*, 5(7):416–423, 2022.
- [19] Peter Clark, Isaac Cowhey, Oren Etzioni, Tushar Khot, Ashish Sabharwal, Carissa Schoenick, and Oyvind Tafjord. Think you have Solved Question Answering? Try ARC, the AI2 Reasoning Challenge. *arXiv preprint arXiv:1803.05457*, 2018.
- [20] Karl Cobbe, Christopher Hesse, Jacob Hilton, and John Schulman. Leveraging procedural generation to benchmark reinforcement learning. In *Proceedings of the 37th International Conference on Machine Learning, ICML’20*. JMLR.org, 2020.
- [21] Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. Training Verifiers to Solve Math Word Problems. *arXiv preprint arXiv:2110.14168*, 2021.
- [22] Lee J Cronbach and Paul E Meehl. Construct validity in psychological tests. *Psychological bulletin*, 52(4):281, 1955.
- [23] David Dalrymple, Joar Skalse, Yoshua Bengio, Stuart Russell, Max Tegmark, Sanjit Seshia, Steve Omohundro, Christian Szegedy, Ben Goldhaber, Nora Ammann, et al. Towards Guaranteed Safe AI: A Framework for Ensuring Robust and Reliable AI Systems. *arXiv preprint arXiv:2405.06624*, 2024.
- [24] Niek F de Jonge, Kevin Mildau, David Meijer, Joris JR Louwen, Christoph Bueschl, Florian Huber, and Justin JJ van der Hooft. Good practices and recommendations for using and benchmarking computational metabolomics metabolite annotation tools. *Metabolomics*, 18(12):103, 2022.
- [25] Alex Dekhtyar and Jane Huffman Hayes. Good benchmarks are hard to find: Toward the benchmark for information retrieval applications in software engineering. 2006.
- [26] Jwala Dhamala, Tony Sun, Varun Kumar, Satyapriya Krishna, Yada Pruksachatkun, Kai-Wei Chang, and Rahul Gupta. BOLD: Dataset and Metrics for Measuring Biases in Open-Ended Language Generation. In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency, FAccT ’21*, page 862872, New York, NY, USA, 2021. Association for Computing Machinery.
- [27] Michael Feffer, Anusha Sinha, Wesley Hanwen Deng, Zachary C. Lipton, and Hoda Heidari. Red-Teaming for Generative AI: Silver Bullet or Security Theater? *arXiv preprint arXiv:2401.15897*, 2024.
- [28] GitHub. Adding a workflow status badge. <https://docs.github.com/en/actions/monitoring-and-troubleshooting-workflows/adding-a-workflow-status-badge>, 2024.

- 453 [29] Shahriar Golchin and Mihai Surdeanu. Time Travel in LLMs: Tracing Data Contamination in
454 Large Language Models. *CoRR*, abs/2308.08493, 2023.
- 455 [30] Declan Grabb, Max Lamparth, and Nina Vasan. Risks from Language Models for Automated
456 Mental Healthcare: Ethics and Structure for Implementation. In *First Conference on Language*
457 *Modeling*, 2024.
- 458 [31] Caglar Gulcehre, Ziyu Wang, Alexander Novikov, Tom Le Paine, Sergio Gómez Colmenarejo,
459 Konrad Zolna, Rishabh Agarwal, Josh Merel, Daniel Mankowitz, Cosmin Paduraru, Gabriel
460 Dulac-Arnold, Jerry Li, Mohammad Norouzi, Matt Hoffman, Nicolas Heess, and Nando de
461 Freitas. RL unplugged: a suite of benchmarks for offline reinforcement learning. In *Proceedings*
462 *of the 34th International Conference on Neural Information Processing Systems*, NIPS '20, Red
463 Hook, NY, USA, 2020. Curran Associates Inc.
- 464 [32] Muhammad Usman Hadi, al tashi, Rizwan Qureshi, Abbas Shah, Muhammad Irfan, Anas Zafar,
465 Muhammad Bilal Shaikh, Naveed Akhtar, Jia Wu, Seyedali Mirjalili, Qasem Al-Tashi, Amgad
466 Muneer, Mohammed Ali Al-garadi, Gru Cnn, and T5 RoBERTa. Large Language Models: A
467 Comprehensive Survey of its Applications, Challenges, Limitations, and Future Prospects.
- 468 [33] Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and
469 Jacob Steinhardt. Measuring Massive Multitask Language Understanding. In *International*
470 *Conference on Learning Representations (ICLR)*, 2021.
- 471 [34] Dan Hendrycks and Thomas Woodside. Devising ML Metrics. *Center for AI Safety*, 2024.
- 472 [35] Matthew Hutson. Artificial intelligence faces reproducibility crisis. *American Association for*
473 *the Advancement of Science*, 2018.
- 474 [36] Alon Jacovi, Avi Caciularu, Omer Goldman, and Yoav Goldberg. Stop uploading test data
475 in plain text: Practical strategies for mitigating data contamination by evaluation benchmarks.
476 *arXiv preprint arXiv:2305.10160*, 2023.
- 477 [37] Minhao Jiang, Ken Ziyu Liu, Ming Zhong, Rylan Schaeffer, Siru Ouyang, Jiawei Han, and
478 Sanmi Koyejo. Investigating Data Contamination for Pre-training Language Models. *ArXiv*,
479 abs/2401.06059, 2024.
- 480 [38] Minqi Jiang, Jelena Luketina, Nantas Nardelli, Pasquale Minervini, Philip H.S. Torr, Shimon
481 Whiteson, and Tim Rocktäschel. WordCraft: An Environment for Benchmarking Commonsense
482 Agents. In *Workshop on Language in Reinforcement Learning (LaRel)*, 2020.
- 483 [39] Kevin B Johnson, Wei-Qi Wei, Dilhan Weeraratne, Mark E Frisse, Karl Misulis, Kyu Rhee,
484 Juan Zhao, and Jane L Snowdon. Precision medicine, AI, and the future of personalized health
485 care. *Clinical and translational science*, 14(1):86–93, 2021.
- 486 [40] Keller Jordan. Calibrated chaos: Variance between runs of neural network training is harmless
487 and inevitable. *arXiv preprint arXiv:2304.01910*, 2023.
- 488 [41] Sayash Kapoor, Emily M. Cantrell, Kenny Peng, Thanh Hien Pham, Christopher A. Bail,
489 Odd Erik Gundersen, Jake M. Hofman, Jessica Hullman, Michael A. Lones, Momin M. Ma-
490 lik, Priyanka Nanayakkara, Russell A. Poldrack, Inioluwa Deborah Raji, Michael Roberts,
491 Matthew J. Salganik, Marta Serra-Garcia, Brandon M. Stewart, Gilles Vandewiele, and Arvind
492 Narayanan. REFORMS: Consensus-based Recommendations for Machine-learning-based
493 Science. *Science Advances*, 10(18):eadk3452, 2024.
- 494 [42] Sayash Kapoor, Peter Henderson, and Arvind Narayanan. Promises and pitfalls of artificial
495 intelligence for legal applications. *arXiv preprint arXiv:2402.01656*, 2024.
- 496 [43] Sayash Kapoor and Arvind Narayanan. Leakage and the reproducibility crisis in machine-
497 learning-based science. *Patterns*, 4(9):100804, 2023.

- [44] Vasileia Karasavva and Aalia Noorbhai. The real threat of deepfake pornography: A review of canadian policy. *Cyberpsychology, Behavior, and Social Networking*, 24(3):203–209, 2021.
- [45] Max Lamparath, Anthony Corso, Jacob Ganz, Oriana Skylar Mastro, Jacquelyn Schneider, and Harold Trinkunas. Human vs. Machine: Behavioral Differences Between Expert Humans and Language Models in Wargame Simulations. *arXiv preprint arXiv:2403.03407*, 2024.
- [46] David J Lewkowicz. The Concept of Ecological Validity: What Are Its Limitations and Is It Bad to Be Invalid? *Infancy*, 2(4):437–450, 2001.
- [47] Yucheng Li. An Open Source Data Contamination Report for Large Language Models. 2023.
- [48] Percy Liang, Rishi Bommasani, Tony Lee, Dimitris Tsipras, Dilara Soylu, Michihiro Yasunaga, Yian Zhang, Deepak Narayanan, Yuhuai Wu, Ananya Kumar, Benjamin Newman, Binhang Yuan, Bobby Yan, Ce Zhang, Christian Alexander Cosgrove, Christopher D Manning, Christopher Re, Diana Acosta-Navas, Drew Arad Hudson, Eric Zelikman, Esin Durmus, Faisal Ladhak, Frieda Rong, Hongyu Ren, Huaxiu Yao, Jue WANG, Keshav Santhanam, Laurel Orr, Lucia Zheng, Mert Yuksekogun, Mirac Suzgun, Nathan Kim, Neel Guha, Niladri S. Chatterji, Omar Khattab, Peter Henderson, Qian Huang, Ryan Andrew Chi, Sang Michael Xie, Shibani Santurkar, Surya Ganguli, Tatsunori Hashimoto, Thomas Icard, Tianyi Zhang, Vishrav Chaudhary, William Wang, Xuechen Li, Yifan Mai, Yuhui Zhang, and Yuta Koreeda. Holistic Evaluation of Language Models. *Transactions on Machine Learning Research*, 2023. Featured Certification, Expert Certification.
- [49] Q Vera Liao and Ziang Xiao. Rethinking Model Evaluation as Narrowing The Socio-Technical Gap. In *Workshop on Artificial Intelligence & Human Computer Interaction*, volume 1, 2024.
- [50] Stephanie Lin, Jacob Hilton, and Owain Evans. TruthfulQA: Measuring How Models Mimic Human Falsehoods. In Smaranda Muresan, Preslav Nakov, and Aline Villavicencio, editors, *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 3214–3252, Dublin, Ireland, May 2022. Association for Computational Linguistics.
- [51] Xiao Liu, Hao Yu, Hanchen Zhang, Yifan Xu, Xuanyu Lei, Hanyu Lai, Yu Gu, Hangliang Ding, Kaiwen Men, Kejuan Yang, Shudan Zhang, Xiang Deng, Aohan Zeng, Zhengxiao Du, Chenhui Zhang, Sheng Shen, Tianjun Zhang, Yu Su, Huan Sun, Minlie Huang, Yuxiao Dong, and Jie Tang. AgentBench: Evaluating LLMs as Agents. In *The Twelfth International Conference on Learning Representations*, 2024.
- [52] Xiao Liu, Hao Yu, Hanchen Zhang, Yifan Xu, Xuanyu Lei, Hanyu Lai, Yu Gu, Hangliang Ding, Kaiwen Men, Kejuan Yang, Shudan Zhang, Xiang Deng, Aohan Zeng, Zhengxiao Du, Chenhui Zhang, Sheng Shen, Tianjun Zhang, Yu Su, Huan Sun, Minlie Huang, Yuxiao Dong, and Jie Tang. Introducing agentbench v0.2. *Github*, 2024.
- [53] Xiao-Yang Liu, Ziyi Xia, Jingyang Rui, Jiechao Gao, Hongyang Yang, Ming Zhu, Christina Dan Wang, Zhaoran Wang, and Jian Guo. FinRL-Meta: Market Environments and Benchmarks for Data-Driven Financial Reinforcement Learning. In *Thirty-sixth Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2022.
- [54] Yu Lu Liu, Su Lin Blodgett, Jackie Chi Kit Cheung, Q Vera Liao, Alexandra Olteanu, and Ziang Xiao. ECBD: Evidence-Centered Benchmark Design for NLP. *arXiv preprint arXiv:2406.08723*, 2024.
- [55] Manikanta Loya, Divya Anand Sinha, and Richard Futrell. Exploring the Sensitivity of LLMs’ Decision-Making Capabilities: Insights from Prompt Variation and Hyperparameters. *arXiv preprint arXiv:2312.17476*, 2023.

- [56] MingYu Lu, Zachary Shahn, Daby Sow, Finale Doshi-Velez, and H Lehman Li-wei. Is deep reinforcement learning ready for practical applications in healthcare? a sensitivity analysis of duel-ddqn for hemodynamic management in sepsis patients. In *AMIA Annual Symposium Proceedings*, volume 2020, page 773. American Medical Informatics Association, 2020.
- [57] Kit-Kay Mak, Yi-Hang Wong, and Mallikarjuna Rao Pichika. Artificial intelligence in drug discovery and development. *Drug Discovery and Evaluation: Safety and Pharmacokinetic Assays*, pages 1–38, 2023.
- [58] Nestor Maslej, Loredana Fattorini, Raymond Perrault, Vanessa Parli, Anka Reuel, Erik Brynjolfsson, John Etchemendy, Katrina Ligett, Terah Lyons, James Manyika, Juan Carlos Niebles, Yoav Shoham, Russell Wald, and Jack Clark. The AI Index 2024 Annual Report. *Institute for Human-Centered AI*, 2024.
- [59] Timothy R McIntosh, Teo Susnjak, Tong Liu, Paul Watters, and Malka N Halgamuge. Inadequacies of large language model benchmarks in the era of generative artificial intelligence. *arXiv preprint arXiv:2402.09880*, 2024.
- [60] Robert J Mislevy, Linda S Steinberg, and Russell G Almond. Focus article: On the structure of educational assessments. *Measurement: Interdisciplinary research and perspectives*, 1(1):3–62, 2003.
- [61] Arvind Narayanan and Sayash Kapoor. Evaluating LLMs is a minefield. https://www.cs.princeton.edu/~arvindn/talks/evaluating_llms_minefield/, 2023.
- [62] Simon Ott, Adriano Barbosa-Silva, Kathrin Blagec, Jan Brauner, and Matthias Samwald. Mapping global dynamics of benchmark creation and saturation in artificial intelligence. *Nature Communications*, 13(1):6793.
- [63] Alexander Pan, Jun Shern Chan, Andy Zou, Nathaniel Li, Steven Basart, Thomas Woodside, Hanlin Zhang, Scott Emmons, and Dan Hendrycks. Do the rewards justify the means? measuring trade-offs between rewards and ethical behavior in the MACHIAVELLI benchmark. In *Proceedings of the 40th International Conference on Machine Learning, ICML’23*. JMLR.org, 2023.
- [64] Alicia Parrish, Angelica Chen, Nikita Nangia, Vishakh Padmakumar, Jason Phang, Jana Thompson, Phu Mon Htut, and Samuel Bowman. BBQ: A hand-built bias benchmark for question answering. In Smaranda Muresan, Preslav Nakov, and Aline Villavicencio, editors, *Findings of the Association for Computational Linguistics: ACL 2022*, pages 2086–2105, Dublin, Ireland, May 2022. Association for Computational Linguistics.
- [65] Mahima Pushkarna, Andrew Zaldivar, and Oddur Kjartansson. Data cards: Purposeful and transparent dataset documentation for responsible ai. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*, pages 1776–1826, 2022.
- [66] E Fantin Irudaya Raj, M Appadurai, and K Athiappan. Precision farming in modern agriculture. In *Smart Agriculture Automation Using Advanced Technologies: Data Analytics and Machine Learning, Cloud Architecture, Automation and IoT*, pages 61–87. Springer, 2022.
- [67] Deborah Raji, Emily Denton, Emily M. Bender, Alex Hanna, and Amandalynne Paullada. AI and the Everything in the Whole Wide World Benchmark. In J. Vanschoren and S. Yeung, editors, *Advances in Neural Information Processing Systems (NeurIPS)*, 2021.
- [68] David Rein, Betty Li Hou, Asa Cooper Stickland, Jackson Petty, Richard Yuanzhe Pang, Julien Dirani, Julian Michael, and Samuel R Bowman. GPQA: A graduate-level google-proof q&a benchmark. *arXiv preprint arXiv:2311.12022*, 2023.

- [69] Anka Reuel, Lisa Soder, Benjamin Bucknall, and Trond Arne Undheim. Position: Technical Research and Talent is Needed for Effective AI Governance. *Forty-first International Conference on Machine Learning*, 2024.
- [70] Juan-Pablo Rivera, Gabriel Mukobi, Anka Reuel, Max Lamparth, Chandler Smith, and Jacquelyn Schneider. Escalation risks from language models in military and diplomatic decision-making. In *The 2024 ACM Conference on Fairness, Accountability, and Transparency*, pages 836–898, 2024.
- [71] Keisuke Sakaguchi, Ronan Le Bras, Chandra Bhagavatula, and Yejin Choi. WinoGrande: an adversarial winograd schema challenge at scale. *Commun. ACM*, 64(9):99106, aug 2021.
- [72] Melanie Sclar, Yejin Choi, Yulia Tsvetkov, and Alane Suhr. Quantifying Language Models’ Sensitivity to Spurious Features in Prompt Design or: How I learned to start worrying about prompt formatting. *arXiv preprint arXiv:2310.11324*, 2023.
- [73] Atsushi Shirafuji, Takumi Ito, Makoto Morishita, Yuki Nakamura, Yusuke Oda, Jun Suzuki, and Yutaka Watanobe. Prompt sensitivity of language model for solving programming problems. In *New Trends in Intelligent Software Methodologies, Tools and Techniques*, pages 346–359. IOS Press, 2022.
- [74] Aarohi Srivastava, Abhinav Rastogi, Abhishek Rao, et al. Beyond the Imitation Game: Quantifying and extrapolating the capabilities of language models. *Transactions on Machine Learning Research*, 2023.
- [75] Bernd Carsten Stahl and David Wright. Ethics and privacy in AI and big data: Implementing responsible research and innovation. *IEEE Security & Privacy*, 16(3):26–33, 2018.
- [76] Arjun Subramonian, Xingdi Yuan, Hal Daumé III, and Su Lin Blodgett. It takes two to tango: Navigating conceptualizations of NLP tasks and measurements of performance. *arXiv preprint arXiv:2305.09022*, 2023.
- [77] Johannes Söding and Michael Remmert. Protein sequence comparison and fold recognition: progress and good-practice benchmarking. *Current Opinion in Structural Biology*, 21(3):404–411, 2011.
- [78] Makoto Takamoto, Timothy Praditia, Raphael Leiteritz, Dan MacKinlay, Francesco Alesiani, Dirk Pflüger, and Mathias Niepert. PDEBench: An Extensive Benchmark for Scientific Machine Learning. In *Thirty-sixth Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2022.
- [79] Rohan Taori, Ishaan Gulrajani, Tianyi Zhang, Yann Dubois, Xuechen Li, Carlos Guestrin, Percy Liang, and Tatsunori B. Hashimoto. Stanford Alpaca: An Instruction-following LLaMA model. https://github.com/tatsu-lab/stanford_alpaca, 2023.
- [80] Gemini Team, Rohan Anil, Sebastian Borgeaud, Yonghui Wu, Jean-Baptiste Alayrac, Jiahui Yu, Radu Soricut, Johan Schalkwyk, Andrew M Dai, Anja Hauth, et al. Gemini: a family of highly capable multimodal models. *arXiv preprint arXiv:2312.11805*, 2023.
- [81] UK AI Safety Institute. Inspect. https://ukgovernmentbeis.github.io/inspect_ai/, 2024.
- [82] Bertie Vidgen, Adarsh Agrawal, Ahmed M. Ahmed, Victor Akinwande, Namir Al-Nuaimi, Najla Alfaraj, Elie Alhajjar, Lora Aroyo, Trupti Bavalatti, Max Bartolo, Borhane Blili-Hamelin, Kurt Bollacker, Rishi Bomassani, Marisa Ferrara Boston, Siméon Campos, Kal Chakra, Canyu Chen, Cody Coleman, Zacharie Delpierre Coudert, Leon Derczynski, Debojyoti Dutta, Ian Eisenberg, James Ezick, Heather Frase, Brian Fuller, Ram Gandikota, Agasthya Gangavarapu, Ananya Gangavarapu, James Gealy, Rajat Ghosh, James Goel, Usman Gohar, Sujata Goswami, Scott A. Hale, Wiebke Hutiri, Joseph Marvin Imperial, Surgan Jandial, Nick Judd, Felix Juefei-Xu,

- 633 Foutse Khomh, Bhavya Kailkhura, Hannah Rose Kirk, Kevin Klyman, Chris Knotz, Michael
634 Kuchnik, Shachi H. Kumar, Srijan Kumar, Chris Lengerich, Bo Li, Zeyi Liao, Eileen Peters
635 Long, Victor Lu, Sarah Luger, Yifan Mai, Priyanka Mary Mammen, Kelvin Manyeki, Sean
636 McGregor, Virendra Mehta, Shafee Mohammed, Emanuel Moss, Lama Nachman, Dinesh Ji-
637 nenhally Naganna, Amin Nikanjam, Besmira Nushi, Luis Oala, Iftach Orr, Alicia Parrish,
638 Cigdem Patlak, William Pietri, Forough Poursabzi-Sangdeh, Eleonora Presani, Fabrizio Puletti,
639 Paul Röttger, Saurav Sahay, Tim Santos, Nino Scherrer, Alice Schoenauer Sebag, Patrick
640 Schramowski, Abolfazl Shahbazi, Vin Sharma, Xudong Shen, Vamsi Sistla, Leonard Tang,
641 Davide Testuggine, Vithursan Thangarasa, Elizabeth Anne Watkins, Rebecca Weiss, Chris
642 Welty, Tyler Wilbers, Adina Williams, Carole-Jean Wu, Poonam Yadav, Xianjun Yang, Yi Zeng,
643 Wenhui Zhang, Fedor Zhdanov, Jiacheng Zhu, Percy Liang, Peter Mattson, and Joaquin Van-
644 schoren. Introducing v0.5 of the AI Safety Benchmark from MLCommons. *arXiv preprint*
645 *arXiv:2404.12241*, 2024.
- 646 [83] Sumeet Wadhvani. AI Benchmarks: Why GenAI Scoreboards Need an Over-
647 haul. [https://www.spiceworks.com/tech/artificial-intelligence/articles/](https://www.spiceworks.com/tech/artificial-intelligence/articles/are-ai-benchmarks-reliable/)
648 [are-ai-benchmarks-reliable/](https://www.spiceworks.com/tech/artificial-intelligence/articles/are-ai-benchmarks-reliable/), 2024. Accessed on June 2, 2024.
- 649 [84] Boxin Wang, Weixin Chen, Hengzhi Pei, Chulin Xie, Mintong Kang, Chenhui Zhang, Chejian
650 Xu, Zidi Xiong, Ritik Dutta, Rylan Schaeffer, Sang T. Truong, Simran Arora, Mantas Mazeika,
651 Dan Hendrycks, Zinan Lin, Yu Cheng, Sanmi Koyejo, Dawn Song, and Bo Li. DecodingTrust:
652 A Comprehensive Assessment of Trustworthiness in GPT Models. In *Advances in Neural*
653 *Information Processing Systems (NeurIPS)*, 2023.
- 654 [85] Lukas M Weber, Wouter Saelens, Robrecht Cannoodt, Charlotte Soneson, Alexander
655 Hapfelmeier, Paul P Gardner, Anne-Laure Boulesteix, Yvan Saeys, and Mark D Robinson.
656 Essential guidelines for computational method benchmarking. *Genome Biology*, 20:1–12, 2019.
- 657 [86] Laura Weidinger, Jonathan Uesato, Maribeth Rauh, Conor Griffin, Po-Sen Huang, John Mellor,
658 Amelia Glaese, Myra Cheng, Borja Balle, Atoosa Kasirzadeh, et al. Taxonomy of risks posed
659 by language models. In *ACM Conference on Fairness, Accountability, and Transparency*, pages
660 214–229, 2022.
- 661 [87] Mark D Wilkinson, Michel Dumontier, IJsbrand Jan Aalbersberg, Gabrielle Appleton, Myles
662 Axton, Arie Baak, Niklas Blomberg, Jan-Willem Boiten, Luiz Bonino da Silva Santos, Philip E
663 Bourne, et al. The FAIR Guiding Principles for scientific data management and stewardship.
664 *Scientific data*, 3(1):1–9, 2016.
- 665 [88] Chejian Xu, Wenhao Ding, Weijie Lyu, Zuxin Liu, Shuai Wang, Yihan He, Hanjiang Hu, Ding
666 Zhao, and Bo Li. SafeBench: A Benchmarking Platform for Safety Evaluation of Autonomous
667 Vehicles. In *Thirty-sixth Conference on Neural Information Processing Systems Datasets and*
668 *Benchmarks Track*, 2022.
- 669 [89] Jiancheng Yang, Rui Shi, Donglai Wei, Zequan Liu, Lin Zhao, Bilian Ke, Hanspeter Pfister, and
670 Bingbing Ni. MedMNIST v2 - A large-scale lightweight benchmark for 2D and 3D biomedical
671 image classification. *Scientific Data*, 10(1):41, 2023.
- 672 [90] Shuo Yang, Wei-Lin Chiang, Lianmin Zheng, Joseph E. Gonzalez, and Ion Stoica. Rethinking
673 Benchmark and Contamination for Language Models with Rephrased Samples, 2023.
- 674 [91] Weihao Yu, Zhengyuan Yang, Linjie Li, Jianfeng Wang, Kevin Lin, Zicheng Liu, Xinchao Wang,
675 and Lijuan Wang. Mm-vet: Evaluating large multimodal models for integrated capabilities.
676 *arXiv preprint arXiv:2308.02490*, 2023.
- 677 [92] Xiang Yue, Yuansheng Ni, Kai Zhang, Tianyu Zheng, Ruoqi Liu, Ge Zhang, Samuel Stevens,
678 Dongfu Jiang, Weiming Ren, Yuxuan Sun, Cong Wei, Botao Yu, Ruibin Yuan, Renliang Sun,
679 Ming Yin, Boyuan Zheng, Zhenzhu Yang, Yibo Liu, Wenhao Huang, Huan Sun, Yu Su, and
680 Wenhui Chen. MMMU: A Massive Multi-discipline Multimodal Understanding and Reasoning
681 Benchmark for Expert AGI. In *Proceedings of CVPR*, 2024.

- [93] Rowan Zellers, Ari Holtzman, Yonatan Bisk, Ali Farhadi, and Yejin Choi. HellaSwag: Can a Machine Really Finish Your Sentence? In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, 2019.
- [94] Kun Zhou, Yutao Zhu, Zhipeng Chen, Wentong Chen, Wayne Xin Zhao, Xu Chen, Yankai Lin, Ji-Rong Wen, and Jiawei Han. Don't Make Your LLM an Evaluation Benchmark Cheater. *ArXiv*, abs/2311.01964, 2023.

NeurIPS Checklist

The checklist follows the references. Please read the checklist guidelines carefully for information on how to answer these questions. For each question, change the default **[TODO]** to **[Yes]**, **[No]**, or **[N/A]**. You are strongly encouraged to include a **justification to your answer**, either by referencing the appropriate section of your paper or providing a brief inline description. For example:

- Did you include the license to the code and datasets? **[Yes]** See Section ??.
- Did you include the license to the code and datasets? **[No]** The code and the data are proprietary.
- Did you include the license to the code and datasets? **[N/A]**

Please do not modify the questions and only use the provided macros for your answers. Note that the Checklist section does not count towards the page limit. In your paper, please delete this instructions block and only keep the Checklist section heading above along with the questions/answers below.

1. For all authors...

- (a) Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope? **[Yes]** *We support all our claims in Sec. 1 in Sec. 6 and App. F.*
- (b) Did you describe the limitations of your work? **[Yes]** *Limitations are described in Sec. 8 and Sec. 9.*
- (c) Did you discuss any potential negative societal impacts of your work? **[Yes]** *The broader impact of our work, including negative implications, is discussed in Sec. 9.*
- (d) Have you read the ethics review guidelines and ensured that your paper conforms to them? **[Yes]** *We conform to all points in the ethics review. For example, we do not work with PII or otherwise sensitive information and any potential negative impacts of our assessment were discussed in Sec. 9.*

2. If you are including theoretical results...

- (a) Did you state the full set of assumptions of all theoretical results? **[N/A]** *Our work does not involve theoretical results.*
- (b) Did you include complete proofs of all theoretical results? **[N/A]** *Our work does not involve theoretical results.*

3. If you ran experiments (e.g. for benchmarks)...

- (a) Did you include the code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL)? **[Yes]** *The code to replicate results will be added as supplementary material and published as part of a GitHub repo upon publication.*
- (b) Did you specify all the training details (e.g., data splits, hyperparameters, how they were chosen)? **[N/A]** *We're not training a model and hence do not include training details. However, we provide all necessary information to replicate the results in our paper as part of the supplementary material.*

- 726 (c) Did you report error bars (e.g., with respect to the random seed after running experi-
727 ments multiple times)? [Yes] *We report statistical significance results for our results,*
728 *where applicable. See Section 6 and Appendix F.*
- 729 (d) Did you include the total amount of compute and the type of resources used (e.g., type
730 of GPUs, internal cluster, or cloud provider)? [N/A] *We did not train or modify a*
731 *model and hence did not use significant compute resources beyond standard laptops.*
- 732 4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets...
- 733 (a) If your work uses existing assets, did you cite the creators? [Yes] *We assess existing*
734 *benchmarks and cite their creators where we mention them.*
- 735 (b) Did you mention the license of the assets? [Yes] *Given that we do not use, distribute*
736 *or modify the benchmarks we assess, we did not mention their license information. We*
737 *release our assessment and results under the CC BY 4.0 license (Sec. 3).*
- 738 (c) Did you include any new assets either in the supplemental material or as a URL? [Yes]
739 *We provide all assessment results as part of this paper in App. F. They will be included*
740 *as part of a repository of benchmark assessments on our website that we will release*
741 *separately to preserve anonymity.*
- 742 (d) Did you discuss whether and how consent was obtained from people whose data you're
743 using/curating? [N/A] *We did not use people's personal data. We base our assessment*
744 *on publicly available information by the respective benchmark developers.*
- 745 (e) Did you discuss whether the data you are using/curating contains personally identifiable
746 information or offensive content? [N/A] *We do not use any PII data and we mentioned*
747 *in the paper that our content is not offensive.*
- 748 5. If you used crowdsourcing or conducted research with human subjects...
- 749 (a) Did you include the full text of instructions given to participants and screenshots, if
750 applicable? [N/A] *We only conducted information-gathering, unstructured interviews*
751 *without explicit instructions to interviewees. There were no formal instructions. How-*
752 *ever, we did show the assessment criteria to interviewees at some point during each*
753 *unstructured interview and asked for their feedback.*
- 754 (b) Did you describe any potential participant risks, with links to Institutional Review
755 Board (IRB) approvals, if applicable? [N/A] *We only conducted information-gathering*
756 *interviews, which do not fall under the category of research with human subjects and*
757 *hence do need an IRB approval.*
- 758 (c) Did you include the estimated hourly wage paid to participants and the total amount
759 spent on participant compensation? [N/A] *The interviews we conducted were only*
760 *done with voluntary participants that were not compensated.*

A Open Challenges in AI Benchmarking

Per the current state of the field, some benchmark issues are not fully addressable by benchmark developer actions and decisions. This section discusses these issues and directs readers, where possible, to resources which cover these open problems in greater depth.

Quick saturation. Rapid advancements in AI have led to the saturation of many benchmarks. Some benchmarks have been saturated within months of their release [58]. Addressing this issue involves evaluating current model performances and assessing whether the concept has already been solved, and determining if the benchmark can be made challenging given state-of-the-art capabilities of the models tested.

Contamination. In Sec. 4.2, we discuss strategies to mitigate data contamination. However, even when fully adhered to, challenges remain. For example, benchmark developers cannot enforce model developers' use of canary strings to avoid training on benchmark data. Preventing data contamination, particularly in models reliant on large amounts of web-scraped data, is a shared responsibility between benchmark and model developers. [90] offers further description of measures that can be taken on the model developer side. This issue is pressing, as contamination has been demonstrated in both FM [29, 37, 47] and non-FM [43, 41]. Future work across stakeholders is needed to effectively mitigate contamination and preserve benchmark validity.

Poor construct validity. Construct validity refers to the degree to which a test or measurement tool accurately measures the construct it intends to measure [22]. [61] outline factors which make construct validity, especially in FM benchmarking, a challenge. They describe certain properties (e.g. factual accuracy) that arise from the interaction between the model and its user population, rather than from the model alone. To combat this, they suggest incorporating ecologically valid⁷ user interactions into the assessment; yet, given the lack of transparency by model developers into actual user interactions, this criteria is difficult to implement for benchmark developers. Alternately, [23] propose that guarantees be made through formal verification, although this approach has not yet been tested in practice.

Standardization of benchmark reporting. Due to the difficulties with construct validity, most benchmarks cannot provide an absolute signal and instead give a relative one by comparison of models on the same benchmark. This signal is often unavailable to potential model users, as there is no present standardization of benchmark reporting. Model developers report whichever benchmarks they see fit without being obligated to provide a rationale, resulting in inconsistent reporting, especially apparent in the case of benchmarks relating to responsible AI concepts [58]. While this issue does not depend on further research, there is no consensus in theory or practice regarding how benchmark reporting should be standardized. Potential avenues towards standardization include publication of benchmark results through independent entities, market incentives such as government contracts, and mandatory reporting as part of AI legislation.

B Stakeholders

This section details the stakeholders that are involved in benchmark development and use processes.

Benchmark developers. Benchmark developers are the individuals or teams who create benchmarks from scratch (e.g. BIG-Bench [74]), by expanding on previously developed benchmarks (e.g. MedMNIST v2 [89]), by integrating multiple existing benchmarks (e.g. HELM [48]), or by both expanding upon and integrating other benchmarks (e.g. Decoding Trust [84]). This groups objectives are developing benchmarks that accurately and comprehensively assess models capabilities or safety-critical characteristics and establishing standards for AI system evaluations that facilitate comparisons and drive progress on the specified tasks. There are three use cases for benchmark developers of our assessment, checklist, and website:

⁷Ecological validity is the extent to which the findings of a research study are able to be generalized to real-life settings [46]

- They use the checklist to understand best practices and guide their benchmark construction process pre-deployment.
- They use the assessment to score their benchmark after constructing it to understand any shortcomings they may address to improve the overall benchmark quality.
- They can use the website to find related benchmarks and compare their benchmark quality to those.

Model developers. Model developers are the individuals or teams who develop AI models for commercial use (e.g. GPT-4 [3]) or non-commercial purposes (e.g. Alpaca [79]). Their objectives in using benchmarks are demonstrating the performance of their models identifying areas for improvement which can guide model development and to establish credibility and encourage adoption by showcasing favorable relative performance. There are three use case for model developers of our assessment and website:

- They can use the assessment results to decide which benchmarks to report
- Model developers can reference our assessment results in their official reporting to indicate quality differences between benchmarks, if applicable
- Model developers can use our website to find relevant benchmarks to report for their model

Model users. Model users are the individuals, organizations, or businesses which use or modify available AI models for various downstream applications (e.g. a company using ChatGPT to provide customer service). Their objective when using benchmark results is making informed decisions regarding which AI models are most suitable for their specific use cases. There are two use case for model users of our assessment and website:

- If model developers don't reference our or any similar benchmark quality assessment, model users can refer to our assessment results on the website to understand quality differences in benchmarks reported by model developers.
- They can also refer to our benchmark assessment results to decide between two related benchmarks who's results may both be relevant for the model user's application context. If one of these benchmarks has a higher quality, they may decide to prioritize that result based on our assessment.

AI researchers. AI researchers are individuals or teams studying AI and related fields either at non-profits, within academic institutions, in industry, or independently. One of researchers objectives is using benchmarks to evaluate the performance of novel AI architectures, training techniques, and approaches, and to compare these to other systems. Additionally, they have the objective of setting research agendas based on the model limitations and weaknesses revealed by benchmarks. There are two use case for AI researchers of our assessment and website:

- Based on our website and assessment results, AI researchers may analyze benchmarking practices in more detail to understand challenges of benchmark developers and drive research on open questions in AI evaluations and AI benchmarking more broadly.
- They can use our website to understand the overall AI benchmark landscape.

Regulators and standard-setting organizations. Regulators and standard-setting organizations may be affiliated with government agencies, international bodies, and industry associations. In these roles, they are responsible for creating and enforcing standards and regulations for AI development and use. Examples of such entities are the AI Safety Institutes, the ISO, and the EU Commission. The objective of these stakeholders is using benchmarks to assess the compliance of AI models with established regulations, guidelines and standards for traits such as performance, fairness, and safety. For example, the UK AI Safety Institute recently released their *Inspect* evaluation framework [81] that includes several benchmarks that we scored in our assessment, among other evaluation strategies. There are two use case for model users of our assessment and website:

- 854 • Regulators and standard-setting organizations can refer to our checklist to design regula-
855 tory requirements, e.g., by only accepting benchmarks as proof for compliance by model
856 developers that completed certain or all criteria in our checklist
- 857 • They can also mandate that only benchmarks that achieved a certain score on our assessment
858 may be used to proof compliance with regulatory requirements.

859 C Benchmark Lifecycle

860 **Design.** During the design stage, a benchmarks purpose, scope, and structure are defined. This
861 requires developers to identify key aspects of an AI system that the benchmark will assess. Based on
862 this decision, they must determine the tasks, datasets, and evaluation metrics which will be used in
863 their benchmark. To inform these decisions, developers consider the requirements of potential users,
864 possibly collaborating with and gathering feedback from these and other stakeholders.

865 **Implementation.** At this stage, the benchmark is constructed and all necessary components are
866 aggregated. Developers collect, process, and (if applicable) annotate the datasets to be used for their
867 tasks. They then create the evaluation scripts which allow models performance on this data to be
868 measured. So that new models can be evaluated, developers may implement user interfaces and APIs
869 which enable access to and interaction with the benchmark. This stage concludes with the initial
870 testing and validation of benchmark components.

871 **Documentation.** To facilitate the benchmarks use and interpretation, benchmark developers need
872 to create comprehensive documentation. This includes preparing detailed descriptions of benchmark
873 tasks, datasets, and evaluation metrics. Additionally, developers may provide instructions for how to
874 access, use, and submit to the benchmark. Documenting design decisions, limitations, and potential
875 biases enables stakeholders to make informed decisions regarding benchmark use. Creating resources
876 for running the benchmark, such as quick-start guides, code documentation, and examples or tutorials
877 is an essential step for accessibility.

878 **Maintenance.** Once the benchmark and its documentation are released, developers must conduct
879 regular maintenance to ensure ongoing usability. They may monitor benchmark usage and perfor-
880 mance to identify areas for improvement and track users compliance with release requirements. Other
881 tasks at this stage include addressing issues or bugs and incorporating user feedback into updates.
882 Developers can regularly update documentation and support materials. Additionally, they can assess
883 the continued relevance and utility of the benchmark by monitoring performance on the benchmark
884 and responding to community feedback.

885 **Retirement.** The final phase of a benchmarks lifecycle is retirement. Benchmarks are phased out
886 or replaced when they become saturated (i.e. model performance reaches the benchmark metrics
887 ceiling), the task studied loses relevance, or better alternatives emerge. During retirement, developers
888 communicate their plan to stakeholders and can provide guidance on transitioning to alternatives.
889 They archive benchmark data, code, and documentation. As a benchmark is retired, developers may
890 share insights gained with the AI community. Finally, they should clearly mark the benchmark as
891 “retired” on channels for deployment and platforms publishing its results.

892 D List of Assessed Benchmakrs

893 We evaluate these 16 foundation model benchmarks (alphabetical order):

- 894 • AgentBench [51]
- 895 • ARC Challenge [19]
- 896 • BBQ [64]

- BIG-bench [74]
- BOLD [26]
- Codex HumanEval [17]
- DecodingTrust [84]
- GPQA [68]
- GSM8k [21]
- HellaSwag [93]
- Machiavelli [63]
- MLCommons AI Safety v0.5 [82]
- MMLU [33]
- MMMU [92]
- TruthfulQA [50]
- WinoGrande [71]

We evaluate these 8 non-foundation model benchmarks (alphabetical order):

- ALE [11]
- FinRL-Meta [53]
- MedMNIST v2 [89]
- PDEBench [78]
- Procggen [20]
- RL Unplugged [31]
- SafeBench [88]
- Wordcraft [38]

E Sensitivity Analysis Details

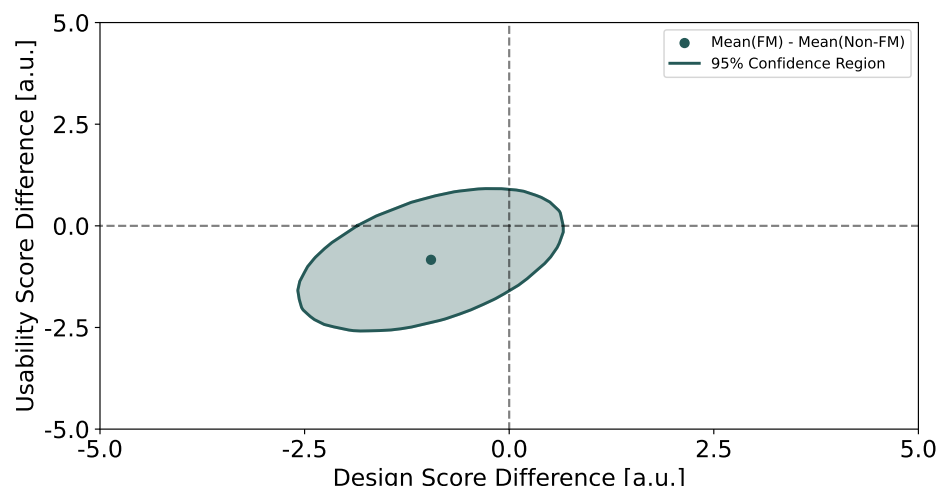


Figure 8: Calculating the difference between the mean Usability and Design score between foundation model (FM) and non-foundation model (Non-FM) benchmarks with the data in Fig. 8. We show the lack of statistical significance of the difference using bootstrap resampling at a 95% confidence level.

We show that the difference in mean usability score between FM and non-FM benchmarks in Fig. 8 is not statistically significant using bootstrap resampling at a 95% confidence level.

922 **F Additional Results**

923 All individual benchmark scoring results, including justifications, can be found on *betterbench.stanford.edu*.
 924

925 **F.1 Scores per lifecycle Stage**

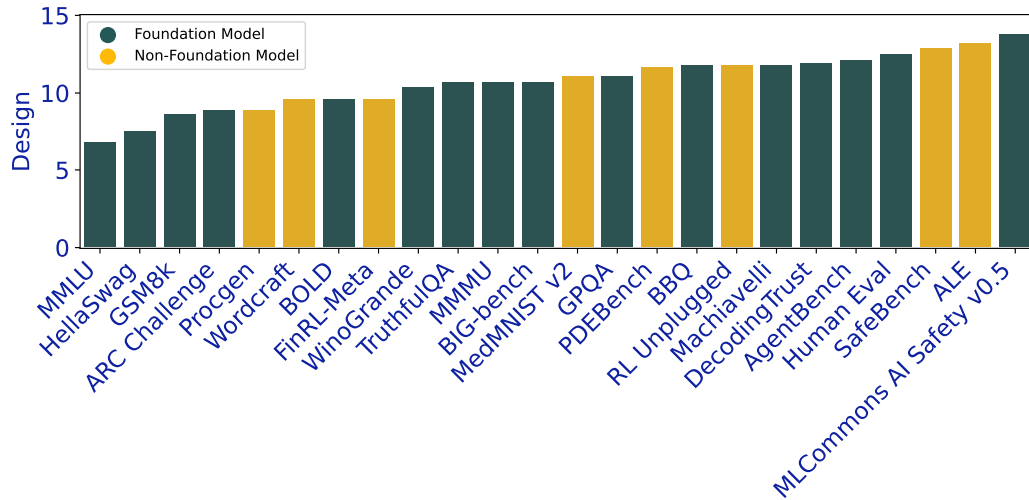


Figure 9: In ascending order, design scores for each benchmark, separated for foundation model (FM) and non-foundation model (Non-FM) benchmarks.

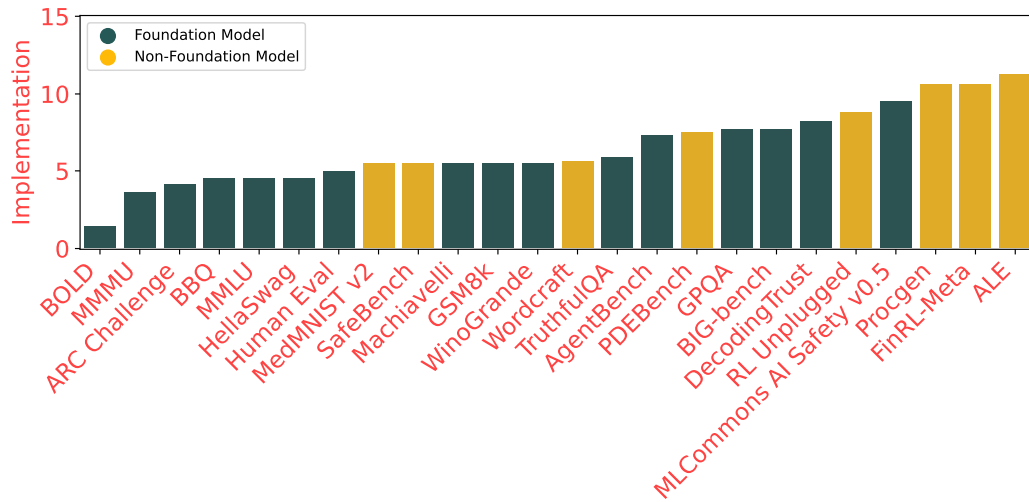


Figure 10: In ascending order, implementation scores for each benchmark, separated for foundation model (FM) and non-foundation model (Non-FM) benchmarks.

926 We show the scores for each benchmark and for each benchmark lifecycle stage as barplots (Design:
 927 Fig. 9, implementation: Fig. 10, documentation: Fig. 11, and maintenance Fig. 12). The scores for
 928 each benchmark for each individual category can be found on our website, *betterbench.stanford.edu*.
 929 For the bar plots for each stage, the benchmarks are shown in ascending order and marked as FM and
 930 non-FM benchmark.

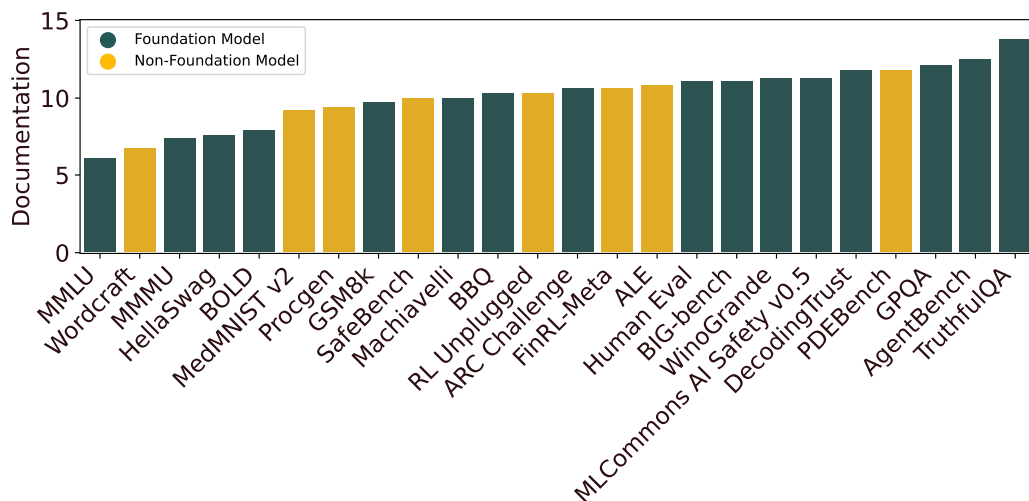


Figure 11: In ascending order, documentation scores for each benchmark, separated for foundation model (FM) and non-foundation model (Non-FM) benchmarks.

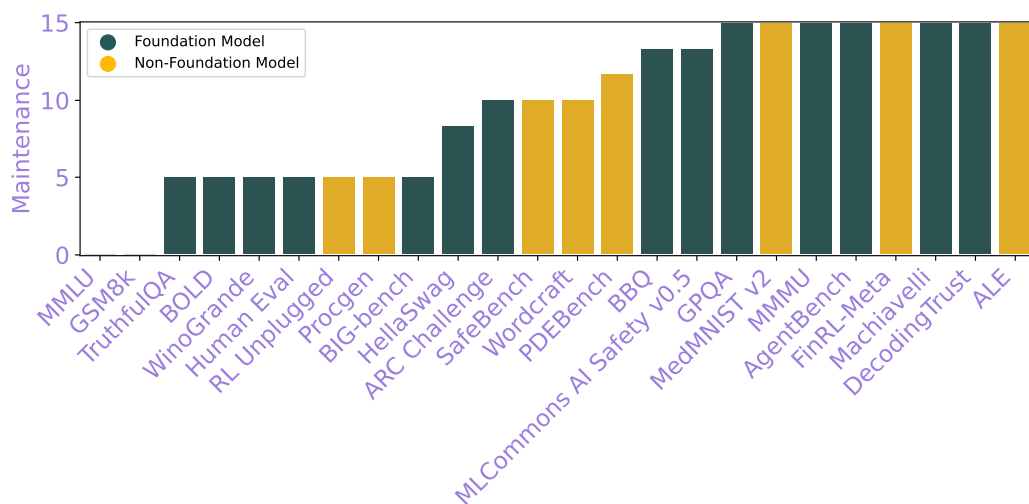


Figure 12: In ascending order, maintenance scores for each benchmark, separated for foundation model (FM) and non-foundation model (Non-FM) benchmarks.

G Scoring

We evaluate 24 benchmarks based on criteria grouped into category (a) (see Sec. 3), i.e., those controlled by the benchmark developer where the authors and interviewees reached a normative consensus. We use the following discrete point system to score each criteria:

- Criteria not acknowledged and not addressed: 0 points
- Criteria acknowledged but not addressed: 5 points
- Criteria partially addressed: 10 points
- Criteria fully addressed: 15 points
- Criteria not relevant: n/a

The highest possible score per category is 15, and the lowest is 0. The criteria span the benchmark lifecycle stages of design, implementation, documentation, and maintenance. Benchmark retirement is excluded from the assessment and scoring, since most benchmarks we looked at are still actively used and not saturated, and given that we cannot predict/anticipate if benchmark developers would in fact fulfill any criteria we'd list for this category. All individual evaluations are made publicly available.

For each lifecycle stage, we calculate the average points earned across the relevant criteria for that stage, excluding any criteria scored as “n/a”. This results in four subscores:

- s_D = Design score
- s_I = Implementation score
- s_{Do} = Documentation score
- s_M = Maintenance score

We do not differentiate the importance of criteria or effort to address them within each lifecycle stage, weighting them equally in the average. To provide an overall assessment of a benchmark's design and usability, we aggregate the subscores into two key measures:

- Design score S_D :
 - Showcases how clear about a benchmark is about its intended purpose and scope and how interpretable it is
 - Equivalent to the design stage subscore s_D
- Usability score S_U :
 - Indicates how easy the benchmark is use and how well it is documented and maintained
 - Weighted average of the implementation, documentation, and maintenance scores, see Equ. 1.

$$S_U = \frac{n_I s_I + n_{Do} s_{Do} + n_M s_M}{n_I + n_{Do} + n_M} \quad (1)$$

Where:

- S_U represents the usability score
- s_I represents the implementation score
- s_{Do} represents the documentation score
- s_M represents the maintenance score
- n_I represents the number of criteria in the implementation stage that are not n/a for the respective benchmark
- n_{Do} represents the number of criteria in the documentation stage that are not n/a for the respective benchmark
- n_M represents the number of criteria in the maintenance stage that are not n/a for the respective benchmark

The discrete 0/5/10/15 point scale provides clearer differentiation between criteria that are not addressed, partially addressed, and fully addressed compared to a continuous scale. At the same time, it allows for a quantitative analysis compared to a letter grade scale like A/B/C/D. Allowing for an N/A option handles criteria that may not be applicable to certain benchmarks. The 0/5/10/15 scale also allows for more granular distinctions compared to a narrower scale like 0/1/2/3 in the final scores: The difference between a score of 5 (acknowledged but not addressed) and 10 (partially addressed) is easier to see than between a 2 and 3 on a narrower scale. With a smaller range, the difference between scores is less meaningful and it is harder to separate the varying degrees of benchmark quality. Providing subscores for each lifecycle stage, while rolling them up into overall Design and Usability Scores, enables assessing benchmarks at both a category and aggregate level.

H Methodology Flow Diagram

Fig. 13 shows a detailed overview of the steps we took to derive the best practices that formed the basis of our AI benchmark assessment.

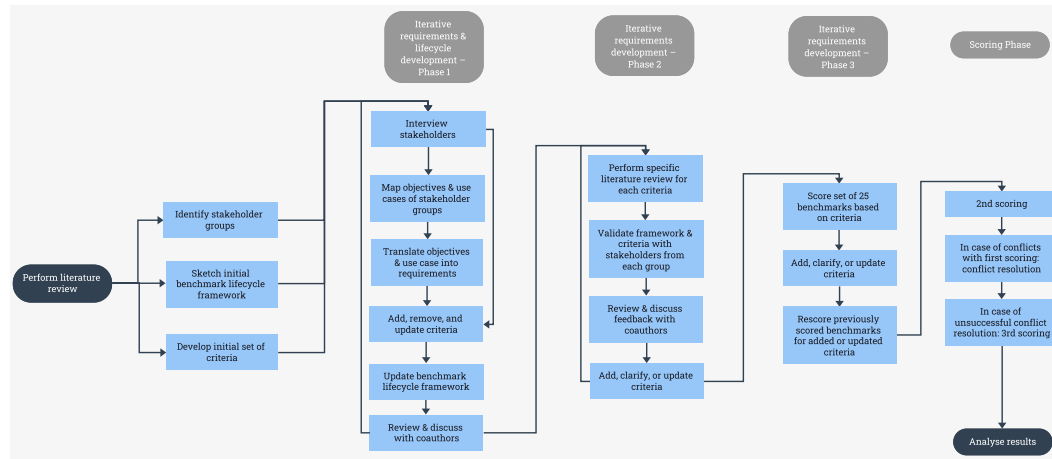


Figure 13: Flow diagram showing our detailed process how we derived the best practices for benchmarks.

I Release Requirements

1. Benchmark developers acknowledge that our checklist is a minimum quality assurance and not sufficient for high-quality benchmark construction.
2. Benchmark developers do not attempt to game our assessment, e.g. by just changing the code checked update on the GitHub repository side without actually checking their code's usability.

J BetterBench Checklist for Benchmark Developers

In this section, we provide the assessment criteria as a checklist for benchmark developers to use during their benchmark construction process, pre-deployment of the benchmark. If benchmark developers want to list their benchmark on our website, they will also have to submit this checklist. On the website, we will further provide an easy-to-fill-out checklist in $\text{\LaTeX}$ and .doc format that can be easily included as part of any benchmark documentation. In the second subsection, we will also add an example of a filled out checklist assessing BetterBench, which can be seen as a benchmark for benchmarks. Going through the checklist was part of the validation of our methodology, described in Step 4 of the Sec. 3 section.

J.1 Template

• Benchmark Design

- ☐ The tested capability, characteristic, or concept is defined
 - TODO | YES | NO | N/A
 - Justification:
- ☐ How tested capability or concept translates to benchmark task is described
 - YES | NO | N/A
 - Justification:
- ☐ How knowing about the tested concept is helpful in the real world is described.

- 1011 – YES | NO | N/A
- 1012 – Justification:
- 1013 ☐ How benchmark score should or shouldn't be interpreted/used is described
- 1014 – YES | NO | N/A
- 1015 – Justification:
- 1016 ☐ Domain experts are involved
- 1017 – YES | NO | N/A
- 1018 – Justification:
- 1019 ☐ Use cases and/or user personas are described
- 1020 – YES | NO | N/A
- 1021 – Justification:
- 1022 ☐ Domain literature is integrated
- 1023 – YES | NO | N/A
- 1024 – Justification:
- 1025 ☐ Informed performance metric choice
- 1026 – YES | NO | N/A
- 1027 – Justification:
- 1028 ☐ Metric floors and ceilings are included
- 1029 – YES | NO | N/A
- 1030 – Justification:
- 1031 ☐ Human performance level is included
- 1032 – YES | NO | N/A
- 1033 – Justification:
- 1034 ☐ Random performance level is included
- 1035 – YES | NO | N/A
- 1036 – Justification:
- 1037 ☐ Automatic evaluation is possible and validated
- 1038 – YES | NO | N/A
- 1039 – Justification:
- 1040 ☐ Differences to related benchmarks are explained
- 1041 – YES | NO | N/A
- 1042 – Justification:
- 1043 ☐ Input sensitivity is addressed
- 1044 – YES | NO | N/A
- 1045 – Justification:
- 1046 • **Benchmark Implementation**
- 1047 ☐ The evaluation code is available
- 1048 – YES | NO | N/A
- 1049 – Justification:
- 1050 ☐ The evaluation data or generation mechanism is accessible
- 1051 – YES | NO | N/A
- 1052 – Justification:
- 1053 ☐ The evaluation of models via API is supported
- 1054 – YES | NO | N/A
- 1055 – Justification:
- 1056 ☐ The evaluation of local models is supported
- 1057 – YES | NO | N/A

- 1058 – Justification:
- 1059 ☐ A globally unique identifier is added or evaluation instances are encrypted
- 1060 – YES | NO | N/A
- 1061 – Justification:
- 1062 ☐ A task to identify if model is included trained on benchmark data
- 1063 – YES | NO | N/A
- 1064 – Justification:
- 1065 ☐ A script to replicate results is explicitly included
- 1066 – YES | NO | N/A
- 1067 – Justification:
- 1068 ☐ Statistical significance or uncertainty quantification of benchmark results is reported
- 1069 – YES | NO | N/A
- 1070 – Justification:
- 1071 ☐ Need for warnings for sensitive/harmful content is assessed
- 1072 – YES | NO | N/A
- 1073 – Justification:
- 1074 ☐ A build status (or equivalent) is implemented
- 1075 – YES | NO | N/A
- 1076 – Justification:
- 1077 ☐ Release requirements are specified
- 1078 – YES | NO | N/A
- 1079 – Justification:

1080 • **Benchmark Documentation**

- 1081 ☐ Requirements file or equivalent is available
- 1082 – YES | NO | N/A
- 1083 – Justification:
- 1084 ☐ Quick-start guide or demo is available
- 1085 – YES | NO | N/A
- 1086 – Justification:
- 1087 ☐ In-line code comments are used
- 1088 – YES | NO | N/A
- 1089 – Justification:
- 1090 ☐ Code documentation is available
- 1091 – YES | NO | N/A
- 1092 – Justification:
- 1093 ☐ Accompanying paper is accepted at peer-reviewed venue
- 1094 – YES | NO | N/A
- 1095 – Justification:
- 1096 ☐ Benchmark construction process is documented
- 1097 – YES | NO | N/A
- 1098 – Justification:
- 1099 ☐ Test tasks & rationale are documented
- 1100 – YES | NO | N/A
- 1101 – Justification:
- 1102 ☐ Assumptions of normative properties are documented
- 1103 – YES | NO | N/A
- 1104 – Justification:

- 1105 ☐ Limitations are documented
 - 1106 – YES | NO | N/A
 - 1107 – Justification:
- 1108 ☐ Data collection, test environment design, or prompt design process is documented
 - 1109 – YES | NO | N/A
 - 1110 – Justification:
- 1111 ☐ Evaluation metric is documented
 - 1112 – YES | NO | N/A
 - 1113 – Justification:
- 1114 ☐ Applicable license is specified
 - 1115 – YES | NO | N/A
 - 1116 – Justification:
- 1117 • **Benchmark Maintenance**
 - 1118 ☐ Code usability was checked within the last year
 - 1119 – YES | NO | N/A
 - 1120 – Justification:
 - 1121 ☐ Maintained feedback channel for users is available
 - 1122 – YES | NO | N/A
 - 1123 – Justification:
 - 1124 ☐ Contact person is listed
 - 1125 – YES | NO | N/A
 - 1126 – Justification:

1127 J.2 Example

1128 As noted in Sec. 3, we assessed BetterBench against our own assessment framework to verify that the
 1129 framework is usable and practicable. This section showcases this assessment and gives an example of
 1130 a filled-out checklist, based on the template provided in App. J.1,

- 1131 • **Benchmark Design**
 - 1132 ☐ The tested capability, characteristic, or concept is defined
 - 1133 – YES
 - 1134 – Justification: “We define a *high-quality* benchmark to be one that is clear about its
 1135 intended purpose and scope, and that is usable. To date, no structured assessment
 1136 for the quality of AI benchmarks, including both FM and non-FM benchmarks, has
 1137 been published to date, and no comparative analysis was conducted to understand
 1138 quality differences between widely used benchmarks in the field. This paper
 1139 addresses these gaps”(Sec. 1)
 - 1140 ☐ How tested capability or concept translates to benchmark task is described
 - 1141 – YES
 - 1142 – Justification: For detail, see Sec. 4 and App. K
 - 1143 ☐ How knowing about the tested concept is helpful in the real world is described.
 - 1144 – YES
 - 1145 – Justification: Justification: “By releasing the first systematic assessment framework
 1146 for AI benchmarks, we aim to encourage benchmark developers to construct higher-
 1147 quality benchmarks and to contribute to community efforts to make AI evaluations
 1148 more practicable and transparent. Higher-quality benchmarks resulting from the
 1149 adoption of our framework and checklist can lead to better-informed model selection
 1150 for downstream tasks, potentially reducing risks and improving outcomes in high-
 1151 stakes applications” (Sec. 9).

- 1152 ☐ How benchmark score should or shouldn't be interpreted/used is described
- 1153 – YES
- 1154 – Justification: “Our living repository of benchmark assessments promotes transparency and comparability, allowing benchmark users to make informed decisions
- 1155 when choosing benchmarks. However, there is a potential risk of misinterpretation
- 1156 of our results; our assessment only provides minimum quality assurances and is not
- 1157 sufficient to assess the suitability of a benchmark for a concrete use case” (Sec. 9).
- 1158
- 1159 ☐ Domain experts are involved
- 1160 – YES
- 1161 – Justification: “Initially, we surveyed the existing benchmark landscape (Sec. 2).
- 1162 Based on this review, we identified five stakeholder groups who present the user
- 1163 personas of our assessment (App. B). All stakeholder groups were represented
- 1164 in subsequent unstructured interviews which included 20+ policymakers, model
- 1165 developers, benchmark developers, model users, and AI researchers, to understand
- 1166 their objectives w.r.t. benchmarking. During this process, we developed a five-
- 1167 stage model of the benchmark lifecycle (Fig. 5 and App. C) and mapped the
- 1168 benchmarking objectives of the stakeholders, along with their communicated use
- 1169 cases of a benchmark assessment (App. B)” (Sec. 3).
- 1170 ☐ Use cases and/or user personas are described
- 1171 – YES
- 1172 – Justification: “We identified five stakeholder groups who present the user personas
- 1173 of our assessment” (Sec. 3, see full personas and use cases in App. B).
- 1174 ☐ Domain literature is integrated
- 1175 – YES
- 1176 – Justification: “Our work is informed by benchmarking practices from fields beyond AI, ranging from transistor hardware [18] to environmental quality [16] to
- 1177 bioinformatics [7], and identify common themes regarding what constitutes an
- 1178 effective benchmark. When applicable, we incorporate these best practices into
- 1179 our assessment (Sec. 4).” Citations for this literature, when used, are provided in
- 1180 Sec. 4.
- 1181
- 1182 ☐ Informed performance metric choice
- 1183 – YES
- 1184 – Justification: “The discrete 0/5/10/15 point scale provides clearer differentiation
- 1185 between criteria that are not addressed, partially addressed, and fully addressed
- 1186 compared to a continuous scale. At the same time, it allows for a quantitative
- 1187 analysis compared to a letter grade scale like A/B/C/D. Allowing for an N/A option
- 1188 handles criteria that may not be applicable to certain benchmarks.” Full details on
- 1189 our scoring method are available in App. G.
- 1190 ☐ Metric floors and ceilings are included
- 1191 – YES
- 1192 – Justification: “The highest possible score per category is 15, and the lowest is 0”
- 1193 (App. G).
- 1194 ☐ Human performance level is included
- 1195 – N/A
- 1196 – Justification: In our work, we manually evaluate AI benchmarks; a human could
- 1197 not be used as an evaluation target in our context.
- 1198 ☐ Random performance level is included
- 1199 – N/A
- 1200 – Justification: Random generation cannot constitute an AI benchmark.
- 1201 ☐ Automatic evaluation is possible and validated
- 1202 – N/A

- 1203 – Justification: “Given the varying information sources (official websites, papers,
1204 GitHub repositories published by the benchmark developers that we do consult
1205 to assess benchmarks, and given that they do not follow a standard structure, we
1206 manually evaluate all benchmarks” (Sec. 3).
- 1207 ☐ Differences to related benchmarks are explained
- 1208 – YES
- 1209 – Justification: “Unlike prior studies, such as [59] and [49], which focus on identify-
1210 ing the limitations, our approach offers a practical evaluation, empowering develop-
1211 ers to address shortcomings and enhance benchmark quality directly” (Sec. 2.1).
- 1212 ☐ Input sensitivity is addressed
- 1213 – N/A
- 1214 – Justification: Since our benchmark uses human evaluation, we select a single
1215 phrasing for each criterion. As described in Sec. 3 these phrasings were devel-
1216 oped iteratively to maximize clarity and minimize disagreement amongst multiple
1217 annotators of the same benchmark.
- 1218 • **Benchmark Implementation**
- 1219 ☐ The evaluation code is available
- 1220 – N/A
- 1221 – Justification: We performed human evaluation which did not use code.
- 1222 ☐ The evaluation data or generation mechanism is accessible
- 1223 – N/A
- 1224 – Justification: We evaluate benchmarks based on “official websites, papers, GitHub
1225 repositories published by the benchmark developers” (Sec. 3). The availability of
1226 these materials is dependent on benchmark developers.
- 1227 ☐ The evaluation of models via API is supported
- 1228 – N/A
- 1229 – Justification: We evaluate benchmarks rather than models.
- 1230 ☐ The evaluation of local models is supported
- 1231 – N/A
- 1232 – Justification: We evaluate benchmarks rather than models.
- 1233 ☐ A globally unique identifier is added or evaluation instances are encrypted
- 1234 – N/A
- 1235 – Justification: Our benchmark does not evaluate AI models or include any examples
1236 which they could be contaminated by training on.
- 1237 ☐ A task to identify if model is included trained on benchmark data
- 1238 – N/A
- 1239 – Justification: Our benchmark does not evaluate AI models or include any examples
1240 which they could be contaminated by training on.
- 1241 ☐ A script to replicate results is explicitly included
- 1242 – N/A
- 1243 – Justification: The code to replicate results will be added as supplementary material
1244 and published as part of a GitHub repo upon publication.
- 1245 ☐ Statistical significance or uncertainty quantification of benchmark results is reported
- 1246 – YES
- 1247 – Justification: These results are reported in Sec. 6 and App. E.
- 1248 ☐ Need for warnings for sensitive/harmful content is assessed
- 1249 – YES
- 1250 – Justification: “The outputs of our evaluation do not contain sensitive or harmful
1251 content, but users may encounter such content during a benchmark assessment
1252 depending on the benchmark’s data” (Sec. 9).

- 1253 ☐ A build status (or equivalent) is implemented
1254 – YES
1255 – Justification: A build status will be included in the code released as part of a GitHub
1256 repo upon publication.
- 1257 ☐ Release requirements are specified
1258 – YES
1259 – Justification: Release requirements are provided in App. I.
- 1260 • **Benchmark Documentation**
- 1261 ☐ Requirements file or equivalent is available
1262 – YES
1263 – Justification: A requirements file will be included in the code released as part of a
1264 GitHub repo upon publication.
- 1265 ☐ Quick-start guide or demo is available
1266 – YES
1267 – Justification: We provide a checklist to facilitate use of our benchmark in App. J
1268 and an example of its use in App. J.2. Additionally, we will include a quick-start
1269 guide for our code in the GitHub repo released upon publication.
- 1270 ☐ In-line code comments are used
1271 – YES
1272 – Justification: Our GitHub repository includes in-line code comments.
- 1273 ☐ Code documentation is available
1274 – YES
1275 – Justification: Our GitHub repository includes code documentation.
- 1276 ☐ Accompanying paper is accepted at peer-reviewed venue
1277 – N/A
1278 – Justification: Our paper is currently under submission at a peer-reviewed venue.
- 1279 ☐ Benchmark construction process is documented
1280 – YES
1281 – Justification: We describe our full process in Sec. 3.
- 1282 ☐ Test tasks & rationale are documented
1283 – YES
1284 – Justification: Definitions and justifications for all criteria are presented in App. K.
- 1285 ☐ Assumptions of normative properties are documented
1286 – YES
1287 – Justification:
- 1288 ☐ Limitations are documented
1289 – YES
1290 – Justification: We discuss limitations in Sec. 8.
- 1291 ☐ Data collection, test environment design, or prompt design process is documented
1292 – YES
1293 – Justification: We describe how we performed our evaluations in Sec. 3.
- 1294 ☐ Evaluation metric is documented
1295 – YES
1296 – Justification: “We define a *high-quality* benchmark to be one that is interpretable
1297 and clear about its intended purpose and scope, and that is usable” Sec. 1. We
1298 further describe how we operationalized “quality” and calculate its subcomponents
1299 (design and usability) in Fig. 9 and Sec. 3.
- 1300 ☐ Applicable license is specified

- 1301 – YES
- 1302 – Justification: We release our assessment under CC BY 4.0 license, available on our
- 1303 website (Sec. 3).
- 1304 • **Benchmark Maintenance**
- 1305 ☐ Code usability was checked within the last year
- 1306 – YES
- 1307 – Justification: We have checked the usability of the code in our GitHub repository
- 1308 and will verify it again upon publication.
- 1309 ☐ Maintained feedback channel for users is available
- 1310 – YES
- 1311 – Justification: “Finally, we develop a supplementary website to continuously publish
- 1312 assessment results using the scoring methodology in App. G, given the rapid
- 1313 development of new benchmarks. The website includes a community feedback
- 1314 channel for submitting new AI benchmarks and correcting previously posted scores
- 1315 if benchmarks are updated or stakeholders disagree with our evaluation” (Sec. 3).
- 1316 ☐ Contact person is listed
- 1317 – YES
- 1318 – Justification: Contact details will be listed on our website.

1319 **K Full Assessment Criteria**

1320 **K.1 Benchmark Design**

1321 **1. Definition of tested capability or characteristic**

- 1322 • **Explanation:** The benchmark developers mention and define what underlying capabil-
- 1323 ity or characteristic of a model is supposed to be tested with the benchmark.
- 1324 • **Justification:** Defining the objective of the benchmark is necessary for clarity in
- 1325 its design. It also helps users determine if the benchmark aligns with their specific
- 1326 application needs and ensures that users and developers have a shared understanding of
- 1327 the concept being evaluated, facilitating consistent interpretation of results.
- 1328 • **Points:**
- 1329 – 0: Tested concept, capability, or characteristic not explicitly mentioned.
- 1330 – 5: Tested concept explicitly mentioned and need for definition acknowledged, but
- 1331 definition not provided.
- 1332 – 10: Tested concept, capability, or characteristic explicitly mentioned but not defined.
- 1333 – 15: Tested concept, capability, or characteristic explicitly mentioned and defined.

1334 **2. Description of how tested capability or concept translates to benchmark task**

- 1335 • **Explanation:** The benchmark developers describe how the tested capability or charac-
- 1336 teristic translates to the task implemented in the benchmark/the task the model is tested
- 1337 on in the benchmark.
- 1338 • **Justification:** Clearly explaining this translation ensures that the benchmark tasks accu-
- 1339 rately reflect the intended tested capabilities and concepts, providing valid assessment
- 1340 results.
- 1341 • **Points:**
- 1342 – 0: No description of how the tested capability or concept translates to the benchmark
- 1343 task.
- 1344 – 5: Acknowledgement that not describing how the tested capability or concept
- 1345 translates to the benchmark task is an issue, but no description provided.
- 1346 – 10: Description of how tested capability or concept translates to benchmark tasks
- 1347 provided for some but not all tasks.

- 1348 – 15: Description of how tested capability or concept translates to benchmark tasks
1349 provided for all tasks.

1350 **3. Description of how knowing about the tested concept is helpful in the real world**

- 1351 • **Explanation:** The developers describe why it is useful to know about the tested
1352 capability in the real world.
- 1353 • **Justification:** This description helps users understand the practical value of the bench-
1354 mark, demonstrating how the tested capability impacts real-world applications and use
1355 cases.
- 1356 • **Points:**
- 1357 – 0: No description of how knowing about the tested concept is helpful in the real
1358 world.
- 1359 – 5: Acknowledgement that not describing how knowing about the tested concept is
1360 helpful in the real world is an issue, but no description provided.
- 1361 – 10: Limited description of how knowing about the tested concept is helpful in the
1362 real world.
- 1363 – 15: Full description of how knowing about the tested concept is helpful in the real
1364 world.

1365 **4. Description of use cases and user personas for the benchmark**

- 1366 • **Explanation:** A use case for an AI benchmark involves specifying a scenario in
1367 which the AI system will be evaluated. This scenario should include the cultural and
1368 geographic context and the type of interactions between humans and models [82], if
1369 applicable. Additionally, user personas should be defined to represent the different
1370 types of users that might interact with the AI system, if applicable. As a concrete
1371 example, [82] states “The use case for the v0.5 Benchmark is an adult chatting to a
1372 general-purpose assistant in English. The cultural and geographic context is Western
1373 Europe & North America. We define a use case as a set of interactions between human
1374 and model to achieve a goal (or goals). [...] For the v0.5 Benchmark, we are focusing on
1375 three personas: (i) a typical adult user; (ii) an adult user intent on malicious activities,
1376 behaving in a technically non-sophisticated way; and (iii) an adult user at risk of harm,
1377 behaving in a technically non-sophisticated way.”
- 1378 • **Justification:** Use cases set the context and scope of the benchmark. User personas
1379 outline an understanding of the different types of interactions the benchmark developers
1380 anticipate the tested AI system to be used in, e.g., ranging from typical users to those
1381 with specific challenges or malicious intent. This approach ensures that the design of
1382 the benchmark is closely related to real-world applications and that it’s effective across
1383 diverse scenarios.
- 1384 • **Points:**
- 1385 – 0: The benchmark does not include any description of use cases or user personas.
- 1386 – 5: The benchmark acknowledges the importance of use cases or user personas but
1387 does not explicitly formulate or describe them.
- 1388 – 10: The benchmark provides a partial description of use cases or user personas.
- 1389 – 15: The benchmark fully describes use cases and user personas, specifying the
1390 cultural and geographic context, types of human-model interactions (if applicable),
1391 and representing different user types that might interact with the AI system (if
1392 applicable).
- 1393 – n/a: For AI systems that do not involve direct human interaction, such as those
1394 used in industrial automation or scientific simulations, defining user personas is not
1395 relevant. However, real-world use cases should still be specified; in more theoretical
1396 benchmarks, this use case might be to advance research.

1397 **5. Involvement of domain experts**

- **Explanation:** Domain expert(s) who have a professional background or research experience in the concept to be tested are either co-authors of the paper, or were involved in the benchmark design process, i.e., the paper makes clear how they obtained the expertise and how that informed the benchmark design.
- **Justification:** Involving domain experts ensures that the benchmark design is informed by deep, specialized knowledge, increasing its validity and relevance. This expertise helps to create tasks that accurately assess the targeted capabilities and align with real-world scenarios.
- **Points:**
 - 0: None of the authors has a background in the benchmark domain and no external experts were consulted during the design process.
 - 5: The benchmark mentions the necessity for in-domain expertise but doesn't specify any further details.
 - 10: The benchmark mentions that domain experts were consulted but not how their insights influenced the benchmark design.
 - 15: At least one of the co-authors has a professional or academic background in the benchmark domain or the benchmark specified how external experts were consulted and how that influenced the design process.

6. Integration of domain literature

- **Explanation:** The developers cite domain literature in the background section and describe how insights from this literature informed the design of their benchmark or cite relevant domain literature in the benchmark design process.
- **Justification:** By consulting domain-specific literature, benchmark developers can ensure that the tasks and evaluation criteria they include are representative and aligned with the current state of knowledge in the field. This literature often contains valuable insights into best practices, established methodologies, and proven approaches for evaluating the tested concept, which can be incorporated into the benchmark design to enhance its reliability.
- **Points:**
 - 0: The benchmark does not reference domain-specific literature.
 - 5: The benchmark mentions the need to integrate domain literature but did not address it in the background section or design process.
 - 10: The benchmark references domain literature in the background or related work section but does not describe how that domain literature informed the benchmark design process.
 - 15: The benchmark references domain literature throughout the paper and describes how that domain literature informed the benchmark design process.

7. Description of how benchmark score should or shouldn't be interpreted/used

- **Explanation:** The benchmark developers provide information about what benchmark users can and cannot take away from the benchmark score.
- **Justification:** Clarifying the interpretation of benchmark scores prevents misuse and misinterpretation, ensuring that users draw accurate conclusions about a model's performance. This guidance helps users apply the scores appropriately within their specific contexts, and understand if the benchmark can be used to assess a model for their desired application context.
- **Points:**
 - 0: The benchmark does not comment on how the benchmark scores should or should not be interpreted.
 - 5: The benchmark acknowledges that the benchmark scores need to be interpreted but gives no guidance on how or how not to do that.

- 10: The benchmark describes how scores should or should not be interpreted or used, but not both.
- 15: The benchmark describes how scores should and should not be interpreted or used.

8. Informed choice of performance metric(s)

- **Explanation:** The developers describe how the performance metric for the defined benchmark task should be interpretable, meaningful, and standard for the task that's being evaluated [34]. If a non-standard metric is selected, they describe their rationale for choosing a non-standard metric.
- **Justification:** The metric should be easily understood by the reader to build their own opinion about the model's capabilities, given the benchmark score. If a non-standard metric is used, an explanation is necessary to clarify its relevance and ensure that users can accurately interpret the results. [34]
- **Points:**
 - 0: The benchmark does not mention an evaluation metric or does not explain the choice of metric.
 - 5: The benchmark acknowledges the need for an informed metric choice but does not justify their metric choice.
 - 10: The benchmark provides an explanation for the choice of some but not all of their metrics.
 - 15: The benchmark provides an explanation for the choice of all of their metrics.

9. Includes floors and ceilings for metric

- **Explanation:** The benchmark provides clear floors and ceilings for the metric(s) it uses [34].
- **Justification:** Establishing clear floors and ceilings for metrics ensures that users have a reference point for understanding model performance. It helps users understand if a benchmark is already saturated or if progress can be made on the task [34]. This also allows benchmark developers to decide when a benchmark should be retired.
- **Points:**
 - 0: The benchmark does not provide any metric floors or ceilings.
 - 5: Floors and ceilings are shown in the results figure but not explicitly mentioned in the text.
 - 10: The benchmark provides floors and ceilings for some but not all evaluation metrics.
 - 15: The benchmark provides floors and ceilings for all evaluation metrics.

10. Includes human performance level

- **Explanation:** The benchmark explicitly states human performance measured on the benchmark task [34]. It also explains how human performance was measured and if this was the performance of an average or expert group of humans. The benchmark notes if measuring human performance is not possible on the benchmark task and why.
- **Justification:** Similar to the previous criteria, including human performance on a benchmark allows the reader to put the model's performance into perspective and allows for a better interpretability of the benchmarking score [34].
- **Points:**
 - 0: The benchmark does not state human performance and does not explain why this is not applicable here.
 - 5: The benchmark mentions human performance in passing but does not provide a measurement or explanation.
 - 10: The benchmark states human performance but does not explain how it was obtained.

- 15: The benchmark states human performance and explains how it was obtained.
- n/a: The benchmark task cannot be completed by a human, and hence reporting human performance is not possible.

11. Includes random performance level

- **Explanation:** The developers explicitly states the random performance measured on the benchmark [34].
- **Justification:** By establishing a baseline performance level achieved through random guessing, generation, or selection, benchmark users can better understand the extent to which a model's performance stems from its inherent capabilities, rather than mere chance or the benchmarks design and especially metric choices. This random performance level serves as a reference point, allowing for a clearer assessment of the model's true effectiveness in tackling the specific task at hand.
- **Points:**
 - 0: The benchmark does not state random performance and does not explain why this is not applicable here.
 - 5: The benchmark mentions random performance but does not provide quantitative random performance on the benchmark task(s).
 - 10: The benchmark states random performance for some but not all tasks.
 - 15: The benchmark states random performance for all tasks.
 - n/a: Measuring random performance on the benchmark task is not possible, and hence reporting random performance is not applicable.

12. Addresses input sensitivity

- **Explanation:** The benchmark contains multiple input variations with the same semantic meaning/intended to elicit the same response or output by the tested model. The developers describe all relevant details such as how many different variations were tested per prompt, and how the variations were designed. For language models, this would mean including a variety of semantically (but not syntactically) equivalent prompts to combat prompt sensitivity [73, 42, 55, 72]. For computer vision models, this could mean inputting a normal, a blurred, and a cropped version of the same image, etc.), while for reinforcement learning, this could mean measuring the sensitivity of learned policies to input features [56].
- **Justification:** Addressing input sensitivity in a benchmark ensures that the model's performance is consistent across semantically equivalent inputs, thus validating its robustness. Including multiple variations per input and detailing their design allows for inspection and replicable evaluation of the model's capabilities. This serves the goal of approximating intrinsic model capabilities or harms better rather than just measuring "an artifact" [61] of your input.
- **Points:**
 - 0: The benchmark does not mention or address input sensitivity.
 - 5: The benchmark mentions the issue of input sensitivity but does not describe experiments to test for it.
 - 10: The benchmark includes some input variations with the same semantic meaning but lacks thorough descriptions or details on the number of variations and their design.
 - 15: The benchmark contains multiple input variations with the same semantic meaning, providing detailed descriptions of all relevant details such as the number of variations per prompt and how they were designed.

13. Validated automatic evaluation available

- **Explanation:** Evaluating a model against a benchmark does not require human evaluation in the process and the quality of the automated evaluation is validated (if applicable, e.g., in the case of FM-based evaluations).

- 1549 • **Justification:** Requiring human feedback to evaluate performance on a benchmark will
1550 significantly limit the scalability of the benchmark and potentially introduce biases from
1551 the human evaluators themselves. In addition, this may require an IRB for researchers,
1552 and will be more costly than an automatic evaluation, creating “major barriers to entry”
1553 [34].
- 1554 • **Points:**
- 1555 – 0: The benchmark does not provide any form of automatic evaluation and relies
1556 entirely on human evaluation.
 - 1557 – 5: The benchmark mentions the benefits of automatic evaluation but provides no or
1558 limited automatic valuation.
 - 1559 – 10: The benchmark includes an automatic evaluation method but does not offer any
1560 validation.
 - 1561 – 15: The benchmark includes an automatic evaluation method and describes how it
1562 was validated as well as the results of the validation.

1563 14. Explanation of differences to related benchmarks

- 1564 • **Explanation:** The benchmark developers explain how their benchmark fills a gap
1565 compared to existing benchmarks or how it expands on existing benchmarks or their
1566 tested concepts.
- 1567 • **Justification:** Benchmark developers demonstrate the added value and relevance of
1568 the new benchmark, justifying its necessity by addressing specific gaps in existing
1569 benchmarks or by expanding on saturated benchmarks. This allows users to better
1570 understand the differences between related benchmarks and determine which one to
1571 use for their specific evaluation context.
- 1572 • **Points:**
- 1573 – 0: The benchmarks do not explain any differences or relevance to existing bench-
1574 marks.
 - 1575 – 5: The benchmark briefly mentions existing benchmarks but provides no explana-
1576 tions of differences or added value.
 - 1577 – 10: The benchmark provides an explanation of how it fills a gap or expands on
1578 existing benchmarks for some but not all mentioned related benchmarks.
 - 1579 – 15: The benchmark provides an explanation of how it fills a gap or expands on
1580 existing benchmarks for all mentioned related benchmarks.

1581 K.2 Benchmark Implementation

1582 1. Availability of evaluation code

- 1583 • **Explanation:** The benchmark developers make the code available for others to evaluate
1584 their own models against the benchmark, e.g., as part of a GitHub repository.
- 1585 • **Justification:**
- 1586 • **Points:** Without access to the benchmarking procedure itself, the benchmark cannot
1587 be scrutinized by external parties to verify its reliability and adequacy, nor can it be
1588 utilized for independent evaluations and comparisons by benchmark users. In addition,
1589 if benchmark users have to write their evaluation code from scratch, its more likely that
1590 seemingly minor implementation details affect the measured performance, hindering a
1591 fair comparison [13].
- 1592 – 0: The evaluation code is not publicly available.
 - 1593 – 5: The benchmark mentions the availability of evaluation code but does not provide
1594 access to it.
 - 1595 – 10: The evaluation code is publicly available for some metrics described by the
1596 benchmark.
 - 1597 – 15: The evaluation code is publicly available for all metrics described by the
1598 benchmark.

2. Script to replicate results is explicitly included

- **Explanation:** The benchmark developers give access to the input, output, and evaluation code, as well as all other necessary information (e.g., hyperparameters or random seed set) that they used to create the initial benchmarking results presented in the paper.
- **Justification:** Providing access to the input, output, and code allows for transparency and reproducibility of the reported results, fostering trust into the benchmark, and contributing to overcome the current reproducibility crisis in AI/ML research [35].
- **Points:**
 - 0: The developers do not provide a script to reproduce the results.
 - 5: The issue of result replicability is mentioned in the benchmark paper but not addressed.
 - 10: A script to reproduce some results in the benchmark paper is available.
 - 15: A script to reproduce all results in the benchmark paper is available.

3. Accessibility of evaluation data, prompts, or dynamic environment

- **Explanation:** The benchmark developers make the evaluation data, prompts, or the data/environment generation mechanism accessible. These do not have to be made public in order to earn full points (if contamination is a concern, for example), but some access to it for evaluation purposes, e.g., by hosting it privately on Hugging Face, needs to be possible.
- **Justification:** Without any accessibility of the evaluation data, prompts, or environment generation mechanism, a benchmark cannot be used.
- **Points:**
 - 0: No access to evaluation data, prompts, or data/environment generation mechanism is provided.
 - 5: The existence of evaluation data, prompts, or data/environment generation mechanism is mentioned, but no concrete access is provided.
 - 10: Partial access to evaluation data, prompts, or data/environment generation mechanism is provided, allowing for limited evaluation.
 - 15: Full access to evaluation data, prompts, or data/environment generation mechanism is provided, enabling comprehensive evaluation.

4. Supports evaluation of models via API calls

- **Explanation:** The benchmark developers allow the benchmark evaluation of models via API access, if applicable.
- **Justification:** This criteria is dependent on the subfield. In NLP, for example, closed-source models such as GPT-4 are oftentimes only accessible via API. Without support for API evaluation, they cannot be evaluated, which is especially problematic if such models are the state-of-the-art models in the field.
- **Points:**
 - 0: The benchmark does not support evaluation of models via API calls.
 - 5: The benchmark mentions the possibility of API evaluation but does not provide concrete implementation details.
 - 10: The benchmark supports evaluation of models via one API.
 - 15: The benchmark supports evaluation of models via two or more APIs to different models.

5. Supports evaluation of local models

- **Explanation:** The benchmark developers implement code to support the evaluation of local models without API access.
- **Justification:** Some model developers only host their models locally. A benchmark should support the evaluation of those to allow for a wide variety of models to be evaluated against the benchmark.

1649 • **Points:**

- 1650 – 0: The benchmark requires users to write their own code to evaluate a local model.
- 1651 – 5: The benchmark mentions that local evaluation should be possible but doesn't
- 1652 provide corresponding code.
- 1653 – 10: The benchmark provides minimal support for local model evaluation, requiring
- 1654 significant user effort.
- 1655 – 15: The benchmark provides full support for local model evaluation with user-
- 1656 friendly code.

1657 **6. Inclusion of a globally unique identifier or encryption of evaluation instances**

- 1658 • **Explanation:** Benchmark developers include a globally unique identifier (GUID) or
- 1659 canary string in the main public evaluation code and all public evaluation prompt or
- 1660 data files. Alternatively, they encrypt the test data files and make the key public.
- 1661 • **Justification:** Including a GUID in relevant (sub-)repositories, public code and data
- 1662 repositories can support the identification of data contamination in models [74], either
- 1663 by allowing model developers to filter out the evaluation data out of large amounts
- 1664 of web-scraped data or by allowing benchmark developers to identify which model
- 1665 developers trained on their data and hence have created models that potentially perform
- 1666 better than they would otherwise on the benchmark. Encrypted test data files prevent
- 1667 non-adversarial crawling of such data; however, [36] advise against “using standard
- 1668 obfuscation or compression methods that are not key-protected, since some crawling
- 1669 systems include pipelines of automatic decompression or deobfuscation.”

1670 • **Points:**

- 1671 – 0: The benchmark does not include a GUID or encryption of evaluation instances.
- 1672 – 5: The benchmark acknowledges the risk of contamination but does not address it.
- 1673 – 10: The benchmark partially implements a GUID or encryption, but not consistently
- 1674 across all relevant files.
- 1675 – 15: The benchmark consistently includes a GUID or encryption across all relevant
- 1676 files and repositories.

1677 **7. Inclusion of 'training_on_test_set' task**

- 1678 • **Explanation:** The benchmark includes a task to identify if the model was trained on
- 1679 the benchmark data.
- 1680 • **Justification:** Public benchmarks face the challenges that their evaluation data may be
- 1681 web-scraped and used to train a model. A 'training_on_test_set' task can serve as a
- 1682 “post-hoc diagnosis of whether [... benchmark] data was used in model training.” [74]
- 1683 • **Points:**
- 1684 – 0: The benchmark does not include a 'training_on_test_set' task.
- 1685 – 5: The benchmark mentions the possibility that models were trained on its data but
- 1686 does not provide a way to check it.
- 1687 – 10: The benchmark includes a partial or limited implementation of a 'train-
- 1688 ing_on_test_set' task that only tests for part of the data used.
- 1689 – 15: The benchmark includes a comprehensive 'training_on_test_set' task.

1690 **8. Assess need for warnings for sensitive/harmful content**

- 1691 • **Explanation:** Benchmark developers explicitly mention in the paper if the evaluation
- 1692 tasks or the expected output may contain sensitive or harmful content. If they do not
- 1693 anticipate sensitive/harmful content in either case, they should explicitly state that.
- 1694 • **Justification:** By explicitly stating the presence of sensitive or harmful content and
- 1695 issuing appropriate warnings, developers help users make informed decisions and take
- 1696 necessary precautions. Even if developers do not expect sensitive or harmful content, if
- 1697 they state that, they showcase to the benchmark users that they actually thought about
- 1698 the possibility. Otherwise, users couldn't be sure if the input or output doesn't contain
- 1699 problematic content or if the developers just forgot to include a warning.

1700 • **Points:**

- 1701 – 0: The benchmark does not mention that they checked for the presence or absence
- 1702 of sensitive/harmful content in the evaluation tasks or expected output.
- 1703 – 5: The benchmark mentions the general possibility of sensitive/harmful content but
- 1704 does not provide clear statements or warnings.
- 1705 – 10: The benchmark explicitly states the presence or absence of sensitive/harmful
- 1706 content for either the evaluation tasks or the expected output.
- 1707 – 15: The benchmark explicitly states the presence or absence of sensitive/harmful
- 1708 content for both the evaluation tasks and the expected output.

1709 9. **Release requirements specified**

1710 • **Explanation:** Benchmark developers specify rules for benchmark users to “ensure

1711 the integrity of test results” [82]. While not all benchmark developers will be able to

1712 enforce the release requirements, they should at least specify them. One example is:

1713 “1. Publishers do not train directly on or against the benchmark dataset and retract any

1714 reported results if and when benchmark data is found to have been in training data. 2.

1715 Techniques that are likely to increase the test performance without a commensurate

1716 increase in safety factor are discouraged and may result in benchmark exclusion. [...]”

1717 [82]

1718 • **Justification:** Written terms of use can help to set expectations and have a foundation

1719 to address subsequent contamination or intentional gamification attempts of the bench-

1720 mark. Potential options they could mention in case of release requirement breaches are,

1721 e.g., “publishing public statements correcting the public record” or “resulting in the

1722 [model] being permanently banned from the benchmark” [82]; however, we will not

1723 assess the enforcement ability or potential listed sanctions as part of this criteria, just

1724 the statement of release requirements.

1725 • **Points:**

- 1726 – 0: The benchmark does not specify any release requirements for benchmark users.
- 1727 – 5: The benchmark briefly mentions the issue of potential gameability or misuse by
- 1728 benchmark users but does not provide specific details.
- 1729 – 10: The benchmark states dos and donts how to use the benchmark but does not
- 1730 specify these as requirements for use.
- 1731 – 15: The benchmark provides a set of release requirements for benchmark users.

1732 10. **Includes *Build Status* or equivalent**

1733 • **Explanation:** A build status is a feature, typically implemented as a GitHub Action,

1734 that indicates whether the most recent build of the benchmark was successful [28]. It

1735 should be implemented for the benchmark’s evaluation code. It verifies that the code is

1736 running correctly after the latest commit.

1737 • **Justification:** A passing build status signifies that the main evaluation code was usable

1738 at the latest commit [28]. Including a build status or equivalent can help to ensure the

1739 reliability and usability of the evaluation code. It allows benchmark users to quickly

1740 determine if the code is functioning as intended, saving time and effort in identifying

1741 potential issues.

1742 • **Points:**

- 1743 – 0: The benchmark neither references nor implements any form of build status or
- 1744 equivalent.
- 1745 – 5: The benchmark mentions the need for working evaluation code but does not
- 1746 implement it in any meaningful way.
- 1747 – 10: The benchmark partially implements a build status or equivalent by providing
- 1748 the information in a less accessible manner.
- 1749 – 15: The benchmark fully implements a build status or equivalent, clearly displaying
- 1750 the status of the most recent build and providing easy access to the information.

1751 **K.3 Benchmark Documentation**

1752 **1. Requirements file available**

- 1753 • **Explanation:** A requirements or environment file, or equivalent is available.
- 1754 • **Justification:** Ease of use is a key criteria for benchmark adoption. Providing a
- 1755 requirements file allows for the quick installation of relevant packages at the correct
- 1756 versions, e.g., within a virtual environment, to use the evaluation code.
- 1757 • **Points:**
 - 1758 – 0: No requirements file or equivalent is provided.
 - 1759 – 5: A requirements file is mentioned but not provided.
 - 1760 – 10: A requirements file is provided but may be missing some dependencies or
 - 1761 versions.
 - 1762 – 15: A complete and accurate requirements file specifying all necessary dependen-
 - 1763 cies and versions is provided.

1764 **2. Quick-start guide or demo code available**

- 1765 • **Explanation:** The benchmark developers make a quick start guide or demo available
- 1766 that walks step-by-step through how the benchmark can be used.
- 1767 • **Justification:** Similar to the criteria above, ease of use is a key criteria for benchmark
- 1768 adoption. Providing a quick-start guide takes away any guesswork on the user side and
- 1769 allows them to directly set up and use the benchmark without spending extra time on
- 1770 setup issues.
- 1771 • **Points:**
 - 1772 – 0: No quick-start guide or demo code is provided.
 - 1773 – 5: A quick-start guide or demo code is mentioned but not provided.
 - 1774 – 10: A quick-start guide or demo code is provided but may be missing some steps or
 - 1775 details.
 - 1776 – 15: A comprehensive, step-by-step quick-start guide or demo code is provided.

1777 **3. Includes informative In-line code comments**

- 1778 • **Explanation:** In-line code comments state the purpose, inputs, outputs, and functional-
- 1779 ity of each code segment in all files relevant for the benchmark evaluation.
- 1780 • **Justification:** In-line documentation of code enhances clarity, understanding, and
- 1781 reproducibility. It facilitates collaboration, maintainability, and makes debugging easier
- 1782 for benchmark developers and users, should that be necessary.
- 1783 • **Points:**
 - 1784 – 0: No in-line code comments are provided.
 - 1785 – 5: In-line code comments are sparse and do not adequately explain the purpose,
 - 1786 inputs, outputs, or functionality of the code.
 - 1787 – 10: Informative in-line code comments are present for most of the code but may be
 - 1788 lacking in detail or clarity for some code segments.
 - 1789 – 15: Comprehensive and informative in-line code comments are provided for all
 - 1790 relevant code segments, clearly explaining their purpose, inputs, outputs, and
 - 1791 functionality.

1792 **4. Code documentation available**

- 1793 • **Explanation:** A full documentation of the repository and code it entails is publicly
- 1794 available. This includes, for example, an overview of the folder structure, the files in
- 1795 the repo, an explanation of functions in the repo.
- 1796 • **Justification:** Detailed documentation of code enhances clarity, understanding, and
- 1797 reproducibility. It facilitates collaboration, maintainability, and makes debugging easier
- 1798 for benchmark developers and users, should that be necessary.
- 1799 • **Points:**

- 1800 – 0: No code documentation is provided.
- 1801 – 5: Code documentation is mentioned but not provided.
- 1802 – 10: Code documentation is minimal or incomplete, lacking important details about
- 1803 the repository structure and functions.
- 1804 – 15: Comprehensive code documentation is provided, including a clear overview
- 1805 of the folder structure, files in the repo, and detailed explanations of all relevant
- 1806 functions.

1807 5. Documentation of test task categories & rationale

- 1808 • **Explanation:** The benchmark developers define the tasks or task categories a model
- 1809 is tested on and describe the rationale for choosing the tasks or task categories. The
- 1810 rationale should explain how these tasks are relevant to the benchmark’s objectives,
- 1811 what they aim to measure, and why they are important for evaluating the concept or
- 1812 capability to be tested.
- 1813 • **Justification:** Documenting test tasks is essential for transparency and for allowing
- 1814 public scrutiny of the benchmark. The rationale provides insight into the selection
- 1815 process, demonstrating that the tasks are not arbitrary but are carefully chosen to reflect
- 1816 real-world applications and user needs. Both help users decide if the benchmark is
- 1817 adequate for their evaluation contexts.
- 1818 • **Points:**
- 1819 – 0: No documentation of test task categories or rationale is provided.
- 1820 – 5: Test task categories are mentioned but they are neither defined in detail and a
- 1821 rationale for their selection is missing or inadequate.
- 1822 – 10: Test task categories are defined, but the rationale for their selection is not
- 1823 provided.
- 1824 – 15: Test task categories are clearly defined, and a comprehensive rationale is
- 1825 provided, explaining their relevance to the benchmark’s objectives, what they
- 1826 measure, and their importance for evaluating the targeted concept or capability.

1827 6. Documentation of assumptions about normative properties

- 1828 • **Explanation:** If the benchmark measures properties that vary across cultural contexts
- 1829 (e.g., politeness), then normative assumptions are explicitly stated. The benchmark
- 1830 developers clearly define the cultural context and values that the benchmark adheres to,
- 1831 explaining how the measured properties are conceptualized and operationalized within
- 1832 the benchmark.
- 1833 • **Justification:** By explicitly stating normative assumptions, the authors provide trans-
- 1834 parency about the cultural framework and values that guide the benchmark’s design
- 1835 and evaluation criteria, which can subsequently ensure cultural sensitivity and mitigate
- 1836 potential biases. It also facilitates informed decision-making for users of benchmarks,
- 1837 specifically for culture-dependent use cases they’re interested in, such as measuring
- 1838 toxicity or bias, for example.
- 1839 • **Points:**
- 1840 – 0: No documentation of normative assumptions is provided, even though the
- 1841 benchmark measures culturally-dependent properties.
- 1842 – 5: The potential influence and importance of cultural context on the benchmark is
- 1843 acknowledged but normative assumptions aren’t stated.
- 1844 – 10: Normative assumptions are stated, but the explanation of how they are concep-
- 1845 tualized and operationalized within the benchmark is incomplete or lacks clarity.
- 1846 – 15: Normative assumptions are explicitly and clearly stated, defining the cultural
- 1847 context and values that the benchmark adheres to, and explaining how the measured
- 1848 properties are conceptualized and operationalized within the benchmark.

1849 7. Documentation of limitations

- **Explanation:** Benchmark developers outline the limitations of the benchmark, including but not limited to the tasks, contexts, and scenarios that are not covered by the evaluation are acknowledged. It's stated which use cases are out-of-scope.
- **Justification:** Documenting a benchmark's limitations is necessary for users to assess its suitability for their specific evaluation needs. By understanding what the benchmark does not cover, users can make informed decisions about whether the benchmark aligns with their goals and whether additional evaluations (either in the form of other benchmarks or private evaluations) may be required to complement the benchmark's results.
- **Points:**
 - 0: No documentation of the benchmark's limitations is provided.
 - 5: Limitations of AI evaluations more broadly are briefly mentioned but without any detail and not applied to the specific benchmark.
 - 10: Either limitations regarding the applicability and use of the benchmark or limitations of the benchmark design are discussed, but not both.
 - 15: Both limitations regarding the applicability and use of the benchmark and limitations of the benchmark design are comprehensively discussed.

8. Documentation of benchmark construction process

- **Explanation:** Benchmark developers give a detailed account of the design process, including the specific decisions made at each lifecycle stage, the rationale behind them, and any trade-offs or compromises (e.g., balancing complexity vs. practicality) considered.
- **Justification:** Documenting the benchmark design process is essential for transparency, as it allows users to understand how the benchmark was created and what factors influenced its development. It allows users to assess the thoroughness and rigor of the benchmark's construction. This information further enables users to critically evaluate whether the benchmark is suitable for their specific use case.
- **Points:**
 - 0: No documentation of the benchmark construction process is provided.
 - 5: The benchmark construction process is briefly mentioned but lacks sufficient detail about the decisions made, rationale, and trade-offs considered.
 - 10: The benchmark construction process is documented, including some decisions made and their rationale, but the description lacks depth or fails to address important aspects such as trade-offs or compromises.
 - 15: The benchmark construction process is comprehensively documented, providing a detailed account of the specific decisions made at each stage, the rationale behind them, and any trade-offs or compromises considered.

9. Provision of a globally unique, persistent identifier for a dataset and its metadata

- **Explanation:** The benchmark dataset and its associated metadata are assigned a globally unique and persistent identifier, such as a Digital Object Identifier (DOI), to ensure long-term accessibility and citability of the resource (FAIR Principles, 2024).
- **Justification:** A persistent identifier supports the findability and accessibility of the benchmark and its dataset. It allows for unambiguous referencing of the data, facilitates proper attribution, and ensures that the dataset can be located and accessed over time, even if its physical location changes. This practice aligns with the FAIR (Findable, Accessible, Interoperable, Reusable) principles, enhancing the benchmark's scientific value and reusability.
- **Points:**
 - 0: The benchmark paper, dataset, and metadata are not assigned any persistent identifier.

- 1900 – 5: The benchmark assigns persistent identifiers to the paper, the dataset, or the
1901 metadata.
- 1902 – 10: The benchmark assigns a persistent identifier to two out of three (paper, dataset,
1903 metadata).
- 1904 – 15: The benchmark assigns a globally unique, persistent identifier to the dataset, its
1905 metadata, and the paper.

1906 10. Inclusion of standardized metadata (e.g., following the Croissant standard)

- 1907 • **Explanation:** The benchmark includes comprehensive, standardized metadata that
1908 describes the dataset, its structure, and relevant information about its creation and usage.
1909 This metadata adheres to established standards such as the Croissant standard, which is
1910 designed specifically for machine learning datasets.
- 1911 • **Justification:** Standardized metadata is crucial for ensuring interoperability and
1912 reusability of the benchmark dataset. It provides consistent and machine-readable
1913 information about the dataset's contents, structure, and provenance. This standard-
1914 ization facilitates easier discovery, understanding, and integration of the dataset into
1915 various research workflows. By following established standards like Croissant, the
1916 benchmark enhances its utility across different platforms and tools in the machine
1917 learning ecosystem.
- 1918 • **Points:**
 - 1919 – 0: The benchmark does not include any structured metadata.
 - 1920 – 5: The benchmark includes some basic metadata, but it is not standardized or
1921 comprehensive.
 - 1922 – 10: The benchmark includes comprehensive metadata that covers most aspects of
1923 the dataset, but it does not fully adhere to a recognized standard like Croissant.
 - 1924 – 15: The benchmark includes complete, standardized metadata (e.g., following the
1925 Croissant standard) that thoroughly describes all aspects of the dataset, ensuring
1926 maximum interoperability and reusability.

1927 11. Documentation of data sources and how the data was collected (if applicable)

- 1928 • **Explanation:** The benchmark provides comprehensive documentation detailing the
1929 origins of the data, the methods used for data collection, and, where applicable, dis-
1930 cusses issues of data provenance and informed consent. They also list the license types
1931 for all data used and how they ensured compliance with that license.
- 1932 • **Justification:** Thorough documentation of data sources and collection methods is
1933 necessary for ensuring transparency, reproducibility, and ethical design of the bench-
1934 mark. It allows users to understand the context and limitations of the data, assess its
1935 appropriateness for their specific use cases, and make informed decisions about its
1936 application. Furthermore, discussing data provenance and informed consent addresses
1937 ethical considerations, particularly when dealing with sensitive or personal data, and
1938 helps ensure compliance with data protection regulations.
- 1939 • **Points:**
 - 1940 – 0: The benchmark provides no information about data sources or collection meth-
1941 ods.
 - 1942 – 5: The benchmark mentions data sources but provides minimal details about
1943 collection methods or ethical considerations.
 - 1944 – 10: The benchmark includes a detailed description of data sources and collection
1945 methods, but lacks a discussion of data provenance, compliance with licensing, or
1946 informed consent, where applicable.
 - 1947 – 15: The benchmark provides extensive documentation of data sources, collection
1948 methods, and a thorough discussion of data provenance, compliance with licensing,
1949 and informed consent, addressing relevant ethical and legal considerations.

1950 12. Documentation of the data preprocessing steps taken

- **Explanation:** The benchmark provides a detailed account of all preprocessing steps applied to the raw data before its inclusion in the final dataset. This documentation includes information on data cleaning, normalization, feature engineering, handling of missing values, and any other transformations or manipulations performed on the original data. If no data preprocessing was done, the authors state this explicitly.
- **Justification:** Thorough documentation of preprocessing steps is necessary for ensuring reproducibility and transparency of the benchmark. It allows users to understand exactly how the final dataset was created, which is key for interpreting results, replicating experiments, and assessing the benchmark's applicability to different use cases. Additionally, this information helps identify potential biases or artifacts introduced during preprocessing that could affect model performance or generalization.
- **Points:**
 - 0: The benchmark provides no information about data preprocessing steps.
 - 5: The benchmark mentions that preprocessing was done but offers minimal details about the specific steps taken.
 - 10: The benchmark includes a general description of preprocessing steps, but lacks comprehensive details or fails to cover all aspects of the data preparation process.
 - 15: The benchmark provides an exhaustive, step-by-step documentation of all preprocessing procedures, including rationales for choices made and potential impacts on the data.

13. Documentation of the data annotation process (if applicable)

- **Explanation:** The benchmark provides documentation of the data annotation process, including the annotation guidelines, the qualifications and training of annotators, the annotation tools used, quality control measures, and inter-annotator agreement metrics. This documentation covers the entire workflow from raw data to the final annotated dataset.
- **Justification:** Comprehensive documentation of the annotation process is necessary for understanding the quality, reliability, and potential biases in the labeled data. It allows users to assess the suitability of the dataset for their specific tasks and to interpret results accurately. Transparent annotation documentation also enables reproducibility of the labeling process, facilitates improvements in future iterations of the benchmark, and helps in identifying and mitigating potential sources of bias or error in the annotations.
- **Points:**
 - 0: The benchmark provides no information about the data annotation process.
 - 5: The benchmark mentions that data was annotated but offers minimal details about the process or guidelines used.
 - 10: The benchmark includes a general description of the annotation process, including guidelines and tools used, but lacks comprehensive details on quality control measures or inter-annotator agreement.
 - 15: The benchmark provides exhaustive documentation of the entire annotation process, including detailed guidelines, annotator information, quality control measures, inter-annotator agreement metrics, and discussions of potential biases or limitations in the annotation approach.

14. Documentation of the representativeness of the data (if applicable)

- **Explanation:** The benchmark provides analysis and documentation of how representative the dataset or environment is of the target population or domain. This includes an explanation of the sampling procedure used, any potential biases in the data collection process, and how well the dataset captures the diversity and distribution of the intended population or phenomenon being studied.
- **Justification:** Understanding the representativeness of the data is necessary for assessing the generalizability and validity of any conclusions drawn from models trained

or evaluated on the benchmark. It helps users identify potential limitations or biases in the dataset that could affect model performance in real-world applications. Proper documentation of representativeness also aids in interpreting benchmark results within the context of the population it represents and highlights areas where the dataset may need expansion or improvement to better cover underrepresented groups or scenarios.

- **Points:**

- 0: The benchmark provides no information about the representativeness of the data or the sampling procedure used.
- 5: The benchmark mentions the importance of data representativeness but offers minimal analysis or explanation of how representative the dataset actually is.
- 10: The benchmark includes a general discussion of data representativeness and the sampling procedure, but lacks comprehensive analysis or fails to address potential biases or limitations in representativeness.
- 15: The benchmark provides an in-depth analysis of data representativeness, including detailed explanation of the sampling procedure, quantitative measures of population coverage, discussion of potential biases, and acknowledgment of any limitations in representativeness.

15. Standardized documentation

- **Explanation:** The benchmark utilizes a standardized documentation format, such as data cards, to present the information about the dataset that is underlying to the benchmark. This standardized approach ensures that all key aspects of the dataset are systematically covered, including its composition, collection methodology, intended uses, ethical considerations, and potential biases.

- **Justification:** Adopting a standardized documentation scheme like data cards enhances the usability and transparency of the benchmark. It provides a consistent, structured format that makes it easier for users to quickly understand the dataset's characteristics, limitations, and appropriate use cases. Standardized documentation facilitates easier comparison between datasets and benchmarks, promotes best practices in data reporting, and helps identify potential issues or gaps in the dataset's coverage.

- **Points:**

- 0: The benchmark does not use any standardized documentation scheme.
- 5: The benchmark includes some elements of standardized documentation, but does not fully adhere to an established scheme like data cards.
- 10: The benchmark uses a standardized documentation scheme, but some sections are incomplete or lack detail.
- 15: The benchmark fully implements a comprehensive standardized documentation scheme (e.g., data cards), providing thorough and structured information on all relevant aspects of the dataset.

16. Documentation of evaluation metric(s)

- **Explanation:** The evaluation metrics used are clearly specified and defined, both for standard and custom metrics tailored to the specific task or domain. The exact formulas or processes used to calculate these metrics, along with any parameters or thresholds employed, are made transparent.

- **Justification:** Documenting the evaluation metrics and scoring process is essential for enabling users to understand how the benchmark quantifies model performance and determines rankings or comparisons. By providing clear and detailed information about the metrics and scoring methods, users can assess whether the chosen metrics are appropriate for the task at hand, align with their own evaluation criteria, and provide a fair and meaningful basis for comparing different models or approaches.

- **Points:**

- 0: No documentation of the evaluation metrics is provided.

- 5: The evaluation metrics are mentioned but not clearly defined, and the exact formulas or processes used to calculate them are not provided.
- 10: The evaluation metrics are defined, but the documentation lacks some important details, such as any parameters or thresholds employed.
- 15: The evaluation metrics are clearly specified. The exact formulas or processes used to calculate these metrics, along with any parameters or thresholds employed, are comprehensively documented.

17. Report statistical significance of benchmark results for at least one model

- **Explanation:** Benchmark developers run statistical significance tests on the benchmark results. They report results for, e.g., more than one random seed, and provide variance bounds. In cases where the benchmark is perfectly deterministic, this should be explicitly stated.
- **Justification:** Not doing statistical significance testing can significantly reduce the validity, utility and confidence in results [13]. Especially for benchmarks, we want to understand how much of the results are due to noise and how much is caused by true differences between the models tested.
- **Points:**
 - 0: No statistical significance testing or variance reporting is provided for the benchmark results.
 - 5: The need for valid benchmarks and/or statistical significance or uncertainty estimation is mentioned but not addressed.
 - 10: Benchmark developers if “bound the expected variation across model training runs” [40], [13]
 - 15: Benchmark developers run statistical significance tests on the benchmark results for at least one model and provide variance bounds or other uncertainty estimations. In cases where the benchmark is perfectly deterministic, this is explicitly stated.

18. Accepted at peer-reviewed venue

- **Explanation:** The benchmark/its associated paper was accepted to a peer-reviewed journal, conference, or similar venue.
- **Justification:** Acceptance at a peer-reviewed venue signifies that the benchmark has undergone an evaluation by an external party, ensuring its validity, reliability, and scientific merit [5]. This peer review process contributes to the credibility and assurance to users that the benchmark meets established standards of quality and relevance [5].
- **Points:**
 - 0: The benchmark/its associated paper has not been accepted at a peer-reviewed venue.
 - 5: The benchmark/its associated paper has been submitted to a peer-reviewed venue but is still under review or awaiting acceptance.
 - 10: The benchmark/its associated paper has been accepted at a peer-reviewed workshop or symposium.
 - 15: The benchmark/its associated paper has been accepted at a peer-reviewed journal, conference, or similar high-profile venue.

19. Specifies applicable license

- **Explanation:** The benchmark developers clearly specify the applicable license for the benchmark in the code repository or paper. This includes providing information about the conditions under which the benchmark can be used, modified, and distributed.
- **Justification:** Specifying the applicable license ensures legal clarity and compliance for benchmark users and enables wider adoption, as commercial users might not be able to use the benchmark if no license is specified.
- **Points:**

- 2103 – 0: No license is specified for the benchmark.
- 2104 – 5: A license is mentioned but not clearly specified or linked to in the code repository
- 2105 or paper.
- 2106 – 10: A license is specified but lacks some important details about the conditions
- 2107 under which the benchmark can be used, modified, or distributed.
- 2108 – 15: The applicable license for the benchmark is clearly specified in the code
- 2109 repository or paper, providing comprehensive information about the conditions
- 2110 under which the benchmark can be used, modified, and distributed.

2111 **K.4 Benchmark Maintenance**

2112 **1. Code usability checked within the last year**

- 2113 • **Explanation:** The main files of the public code were updated within the last year⁸, or
- 2114 the developers checked that the benchmark code is still usable and explicitly state this
- 2115 check in the README file, including the date of the check.
- 2116 • **Justification:** Over time, packages that the benchmark depends on may be updated and
- 2117 become incompatible with the original evaluation/benchmark code. To ensure ongoing
- 2118 usability, benchmark developers must check if their code can still be used at least once
- 2119 a year⁹. This practice ensures that users can use the benchmark without encountering
- 2120 and having to fix issues due to outdated dependencies.
- 2121 • **Points:**
- 2122 – 0: No updates to the main files of the public code within the last year, and no
- 2123 explicit statement of a usability check in the README file.
- 2124 – 5: Updates to minor files in the repo were made (e.g., README file) but an explicit
- 2125 statement of a usability check in the README file is not reported.
- 2126 – 10: Updates to the main files of the public code were made within the last year, but
- 2127 the build status check failed and wasn't fixed.
- 2128 – 15: Updates to the main files of the public code within the last year, accompanied
- 2129 by a successful build status check, or an explicit statement of a usability check in
- 2130 the README file, including the date of the check was provided.

2131 **2. Maintained feedback channel for users**

- 2132 • **Explanation:** GitHub issues are acknowledged or addressed within three months. If
- 2133 there are no open issues, benchmark developers would get full points.
- 2134 • **Justification:** Over time, users may find issues with the benchmark tasks or imple-
- 2135 mentation. To ensure continued usability, benchmark developers should address these
- 2136 concerns in a reasonable amount of time. Promptly responding to user feedback helps
- 2137 maintain the reliability and relevance of the benchmark.
- 2138 • **Points:**
- 2139 – 0: No acknowledgment or response to GitHub issues that are older than three
- 2140 months¹⁰.
- 2141 – 5: GitHub issues are mentioned as a way to provide feedback but there are GitHub
- 2142 issues that were not responded to and that are older than three months.
- 2143 – 10: All GitHub issues are acknowledged within three months, but not all are
- 2144 addressed or resolved or were closed because the issue/feature request won't be
- 2145 attended to.

⁸We recognize that this criterion is just a proxy for checking code usability, but we assume that if the main code was edited and a build status [28] passed, that the usability was sufficiently checked.

⁹The one-year threshold is somewhat arbitrary but out of experience of the authors, there is some transition period until which old versions can still be reliably used and are maintained, which can vary from a few months to a few years.

¹⁰This is an arbitrary cut-off time but it seemed reasonable to give developers extended time to respond to open issues.

2146 – 15: All GitHub issues are acknowledged and addressed within three months, or it is
2147 clearly stated if an issue cannot be fixed or if a feature request won't be fulfilled.
2148 Alternatively, there are no open issues¹¹.

2149 **3. Provide contact details of person responsible for benchmark**

- 2150 • **Explanation:** The benchmark should include contact details of the person responsible,
2151 such as a corresponding author in the associated paper, a contact person listed on
2152 GitHub or the website, or an available online feedback form.
- 2153 • **Justification:** Providing contact details ensures that users have a communication
2154 channel for inquiries, feedback, or reporting issues related to the benchmark. This
2155 transparency supports effective collaboration and resolution of problems, enhancing
2156 the benchmark's usability.
- 2157 • **Points:**
 - 2158 – 0: It is not disclosed who developed the benchmark.
 - 2159 – 5: The benchmark developers are disclosed but no explicit contact details are
2160 provided.
 - 2161 – 10: Contact details are provided but are incomplete or difficult to find, e.g., only as
2162 part of terms of service on a website.
 - 2163 – 15: Contact details of the person responsible for the benchmark are easily accessible,
2164 such as a corresponding author in the associated paper, a contact person listed on
2165 GitHub or the website, or an available online feedback form.

¹¹This is an imperfect proxy for a maintained feedback channel. It may be that the benchmark is working well or it may be that the benchmark is not used enough for issues to occur. However, maintenance is a critical part of benchmarks, and we hence decided to include an imperfect proxy rather than not including this criterion at all.