
Data Centric Guard (DC-Guard) - A Framework for Trustworthy LLM Evaluation

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 In the current era, Large Language Models (LLMs) continue to achieve remarkable
2 results yet their evaluation is increasingly undermined by data-centric challenges
3 such as contamination, memorization and benchmark bias which threaten the
4 reliability of reported performance. To address these issues, we propose DC-
5 Guard (Data Centric Guard), a unified framework for trustworthy evaluation of
6 LLMs. This framework introduces three novel components: the Memorization
7 Consistency Index (MCI) to probe hidden memorization, the Benchmark Ecology
8 Score (BES) to quantify representativeness relative to real-world corpora, and the
9 Contamination-Resilient Metric Adjustment (CRMA) to correct evaluation scores
10 for the contamination risk. Together, these elements provide contamination-aware,
11 bias-adjusted reproducible assessments. Beyond presenting this methodology, we
12 discuss open challenges in maintaining robust evaluations under evolving data
13 sources and shifting usage contexts. DC-Guard offers principled guardrails for fair
14 and transparent benchmarking of the large-scale language models.

15 1 Introduction

16 The conventional benchmarks of LLM evaluation suffer from data contamination due to training-test
17 overlap, hidden memorization of surface patterns and coverage bias arising from narrow benchmark
18 distributions. As a result, the existing evaluation practices risk overestimating true model capabilities
19 and undermining reliability. These data centric challenges raise serious concerns about fairness,
20 reproducibility, and interpretability. The community has responded with a variety of methods such as
21 overlap-based contamination detection, quiz-style memorization probes, ad-hoc measures of dataset
22 bias. Yet, these efforts remain fragmented and non-standardized. Existing studies apply different
23 signals, thresholds and definitions, making it difficult to compare results or establish the common
24 ground.

25 To address this gap, we propose a principled data-centric evaluation framework. The central idea is
26 to treat benchmark evaluation not as a fixed score-reporting exercise, but as an ecological audit of
27 contamination, memorization, and representativeness. This framework unifies scattered techniques
28 into a coherent structure, introducing three new components: the Benchmark Ecology Score (BES) for
29 quantifying coverage bias, the Memorization Consistency Index (MCI) for separating memorization
30 from reasoning, and the Contamination-Resilient Metric Adjustment (CRMA) for reporting fairer
31 accuracy under contamination risk. By reframing evaluation as a structured, contamination-aware
32 ecological process, DC-Guard aims to provide transparent, reproducible, and trustworthy guardrails
33 for both researchers and practitioners.

2 Literature Survey

The central challenges most of the time revolve around data contamination, memorization, and benchmark representativeness. Several strands of work have emerged addressing these namely,

Data Contamination: One of the earliest and most widely discussed challenge where benchmark items or near-duplicates are already present in the pretraining corpus. This compromises the validity of evaluation since the model may recall the material rather than generalizing it. Approaches for detecting contamination generally fall into two categories:

1. **Surface-level checks:** n-gram overlap or token matching methods, which scan for exact or near-exact text overlap between benchmarks and training corpora. While simple, they fail to capture the semantic rephrasings or paraphrases.
2. **Semantic similarity checks:** Embedding-based methods that measure closeness in the representation space, thereby detecting paraphrased or slightly altered duplicates. These approaches are more robust but still lack a clear calibration for what constitutes meaningful contamination.

Recent studies have introduced more structured tools, such as quiz-style probes that evaluate whether a model can distinguish between original benchmark items and perturbed versions. These methods provide stronger evidence of contamination but remain task-specific and often lack standardization.

Memorization: In this issue, even when evaluation items are not directly present in the training data, LLMs reproduce rare or idiosyncratic information memorized during pretraining. This challenges the notion of “generalization,” as models may appear capable when they are in fact recalling. To address this issue, several diagnostic strategies have been proposed such as:

1. **Paraphrase-based Probing:** Checking whether models remain consistent across reworded prompts or equivalent queries. A sharp performance drop under paraphrasing often indicates superficial memorization.
2. **Frequency Analysis:** Investigating whether models disproportionately reproduce sequences that were rare but frequent enough in training to be memorized.
3. **Quiz-style Memorization Detection:** Rephrased or misleading options are introduced to test whether a model is truly reasoning or simply reproducing a memorized pattern.

Despite these advances, memorization detection still struggles with calibration. For instance, if a model answers consistently across paraphrases, is it demonstrating robust reasoning or consistent recall? Current approaches lack a principled metric to separate the two.

Benchmark Representativeness: Another growing concern is the mismatch between benchmarks and real-world usage. Widely used datasets often emphasize narrow domains and stylized problem settings which create a coverage bias where benchmarks may not reflect the diversity of tasks, topics, and linguistic structures encountered in the deployment. Research has sought to quantify representativeness using distributional similarity metrics, comparing benchmarks against large reference corpora. Yet, these studies use different divergence measures and rarely translate their findings into an interpretable score that can be easily adopted with.

In summary, related works have identified the core risks but have addressed them piecemeal. What is missing is a data-centric framework that unifies these strands under a coherent philosophy and provides principled metrics. This motivates our proposal of DC-Guard, which consolidates the fragmented efforts into a single theoretical pipeline.

3 Proposed Framework

We introduce Data Centric Guard (DC-Guard), whose central principle is to reframe benchmark evaluation as an ecological audit. Each benchmark is treated not merely as a static dataset, but as an environment whose validity depends on its freedom from contamination, its resistance to memorization artifacts, and its representativeness of real-world usage.

81 The DC-Guard is organized into three theoretical pillars, each corresponding to a novel contribution.
 82 Benchmark Ecology Score (BES), Memorization Consistency Index (MCI), and Contamination-
 83 Resilient Metric Adjustment (CRMA) are linked in a workflow that begins with auditing the dataset,
 84 proceeds to auditing the model behavior, and finally yields adjusted evaluation scores that more
 85 faithfully represent the true generalization.

86 1. **Benchmark Ecology Score (BES):** Traditional benchmarks are often narrow in scope and
 87 do not reflect the linguistic and topical diversity encountered in the deployment. Hence, BES
 88 quantifies the degree of divergence between a benchmark dataset and a real-world reference
 89 corpus. The score draws inspiration from ecological diversity indices, where ecosystems are
 90 evaluated based on species richness and evenness. Analogously, benchmarks can be viewed
 91 as habitats that sample certain linguistic species (topics, styles, reasoning types).

92 **Computation:**

- 93 (a) Represent each benchmark and reference corpus in a shared distributional space, such
 94 as topic distributions or sentence embeddings.
 95 (b) Compute divergence between the two distributions
 96 (c) Map the divergence to a categorical scale: Low Bias (close alignment), Medium Bias,
 97 or High Bias (substantial mismatch).
 98 2. **Memorization Consistency Index (MCI):** It is designed to disentangle memorization from
 99 reasoning by measuring, how models respond to paraphrased variants of benchmark prompts
 100 while correcting for the background answer frequency.

101 **Workflow:**

- 102 (a) For each benchmark item x , generate paraphrased variants that preserve semantic
 103 meaning but alter surface form.
 104 (b) Compute $Consistency(x)$: the fraction of paraphrases where the model produces
 105 identical answers.
 106 (c) Compute $Background Match(x)$: the probability that same answer arises in unrelated
 107 prompts (capturing rote response patterns).
 108 (d) Define the index:

$$MCI(x) = Consistency(x) \times (1 - BackgroundMatch(x))$$

109 A high MCI typically indicates that responses are driven by memorization rather than
 110 reasoning. Medium MCI suggests mixed signals hence, requires closer inspection. Low
 111 MCI suggests reliance on reasoning or contextual adaptation rather than rote recall.

112 3. **Contamination-Resilient Metric Adjustment (CRMA):** This integrates contamination
 113 detection and memorization into a single adjusted metric, preventing inflation of perfor-
 114 mance scores due to training-test overlap. It modifies raw accuracy scores by discounting
 115 performance proportional to estimated contamination probability, while integrating the
 116 memorization evidence from MCI metric.

117 **Formulation:** Let,

- 118 • Acc = raw accuracy of the model on benchmark items.
 119 • $\hat{C}$ = calibrated contamination probability, derived from null-model similarity distribu-
 120 tions.
 121 • MCI = Memorization Consistency Index for the same items.

122 We define the Adjusted Accuracy as:

$$Acc_{adj} = Acc \times (1 - \hat{C} \cdot f(MCI)) \quad (1)$$

123 where $f(MCI)$ scales contamination penalties by memorization evidence. For instance, if
 124 contamination is detected but memorization is low, the penalty is reduced, thereby avoiding
 125 double penalization. This formulation couples contamination and memorization signals into
 126 a single correction factor, ensuring balanced fairness in evaluation.

4 Discussion

While DC-Guard establishes a structured framework for auditing benchmarks and mitigating contamination and memorization, it also opens up several avenues for deeper inquiry. Primarily, the Benchmark Ecology Score (BES) provides a systematic way of quantifying benchmark representativeness by drawing parallels to ecological diversity. It mainly relies on selecting a reference corpus against which diversity is measured. This inevitably introduces bias, which must be carefully addressed. The ecological alignment must be re-assessed periodically as real-world tasks evolve over time. A major challenge is how do we define and update the ground truth ecology in a dynamic language environment.

The Memorization Consistency Index (MCI) uses paraphrased variants and background-match corrections which bridges the gap between anecdotal evidence of memorization and a quantitative diagnostic tool. Paraphrase-based probing has inherent limitation which is generating faithful paraphrases that preserve difficulty, context, and cultural nuance is nontrivial. Overreliance on these automated paraphrasing tools risks in introducing artifacts. Another challenge lies in differentiating productive memorization like recalling factual constants from unproductive memorization like verbatim recall of benchmark items. A nuanced taxonomy of memorization types needs to be developed.

As LLM applications increasingly shift toward interactive, multimodal, and real-time settings, the static text benchmarks alone may prove inadequate. Extending the framework to multi-turn dialogues, multimodal datasets, and task-specific evaluation remains an open frontier. DC-Guard provides three metrics that could serve as standardized reporting tools, improving comparability across studies. But, Standardization itself is difficult: different research groups may operationalize BES or MCI differently depending on corpora, paraphrasing techniques and contamination baselines. Model developers could optimize them without genuinely improving the generalization. Moreover, publicly flagging contaminated benchmarks may inadvertently disincentivize dataset re-use, even when re-use is valid. Balancing transparency with practicality and aligning evaluation standards with policy frameworks, remain as an unresolved societal challenge. Without community-wide guidelines, the results may diverge. There is a need for shared benchmarks, open-source toolkits, and consensus protocols to ensure reproducibility.

5 Conclusion

LLMs have reached unprecedented levels of fluency and task coverage but their evaluation pipelines remain deeply entangled with data centric challenges. Traditional methods often underemphasize these problems, leading to inflated claims of progress and unreliable signals. In this paper, we introduced DC-Guard, a framework that rethinks evaluation from a data-centric perspective.

By combining three dimensions namely Benchmark Ecology Score (BES) to assess benchmark representativeness, Memorization Consistency Index (MCI) to diagnose model recall behavior, and Contamination-Resilient Metric Adjustment (CRMA) to correct inflated scores, the framework provides a holistic approach for auditing the LLM evaluation pipeline. The strength lies not only in its individual components but also in its integration of ecology, memorization, and contamination into a single evaluative lens. This multidimensional view encourages the field to move beyond and achieve more trustworthy measures of generalization.

References

- [1] Golchin, S. & Surdeanu, M. (2023) Data Contamination Quiz: A Tool to Detect and Estimate Contamination in Large Language Models. *arXiv preprint arXiv:2311.06233*.
- [2] Xu, C., Yan, N., Guan, S., Jin, C., Mei, Y., Guo, Y. & Kechadi, M.-T. (2025) DCR: Quantifying Data Contamination in LLMs Evaluation. *arXiv preprint arXiv:2507.11405*.
- [3] Dong, Y., Jiang, X., Liu, H., Jin, Z., Gu, B., Yang, M. & Li, G. (2024) Generalization or Memorization: Data Contamination and Trustworthy Evaluation for Large Language Models. In *Findings of ACL 2024*, pp. 12039–12050. :contentReference[oaicite:0]index=0
- [4] Xu, C. & Yan, N. (2025) TripleFact: Defending Data Contamination in the Evaluation of LLM-driven Fake News Detection. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Long Papers)*, pp. 8808–8823. :contentReference[oaicite:1]index=1
- [5] Singh, A.K., Kocyigit, M.Y., Poulton, A., Esiobu, D., Lomeli, M., Szilvasy, G. & Hupkes, D. (2024) Evaluation Data Contamination in LLMs: How Do We Measure It and (When) Does It Matter? *arXiv preprint arXiv:2411.03923*.
- [6] Xu, C., Guan, S., Greene, D. & Kechadi, M.-T. (2024) Benchmark Data Contamination of Large Language Models: A Survey. *arXiv preprint arXiv:2406.04244*.
- [7] Palavalli, M., Bertsch, A. & Gormley, M.R. (2024) A Taxonomy for Data Contamination in Large Language Models. In *Proceedings of the 1st Workshop on Data Contamination (CONDA)*, pp. 22–40. :contentReference[oaicite:2]index=2
- [8] Nielsen, A.L. & Jordan, M.I. (2022) Statistical Calibration of Semantic Similarity Scores. *Journal of Machine Learning Research*, **23**(140):1–32.
- [9] Murphy, N.C., Kulkarni, S. & Haas, P.J. (2023) Domain Representativeness in Language Understanding Benchmarks. In *Findings of ACL*, pp. 987–1000.

Supplementary Material

The lifecycle of an LLM typically unfolds across three stages: pretraining, where the model learns statistical regularities from massive text corpora; fine-tuning and alignment, where the model is adapted to follow instructions or domain-specific data; and evaluation, where performance is reported on benchmarks meant to represent real-world use. Research on evaluation of LLMs has increasingly shifted towards identifying the flaws in the standard benchmarking practices.

Let us view the usage of the theoretical definitions of DC-Guard Metrics namely Benchmark Ecology Score (BES), Memorization Consistency Index (MCI) and Contamination-Resilient Metric Adjustment (CRMA) on widely adopted benchmarks:

Example 1

Let us consider evaluating a model on the widely used SQuAD (Stanford Question Answering Dataset) benchmark, with Wikipedia (2023 snapshot) serving as the reference corpus. Suppose the model has been fine-tuned on QA tasks, we are now interested in quantifying how contamination and memorization affect reported accuracy.

Assumptions:

- Raw accuracy (Acc): 85%
- Benchmark Ecology Score (BES): Medium Bias (factual QA coverage, under-represents reasoning/dialogue diversity)
- Contamination probability ($\hat{C}$): 0.25 (25% overlap chance with pretraining corpus)
- Memorization Consistency Index (MCI): 0.70 (high consistency across paraphrases)
- Function $f(MCI) = MCI$ (linear scaling)

212 **Inferences:**

- 213 1. **BES:** Medium Bias $\Rightarrow$ Benchmark is focused but not fully representative; lacks reasoning
214 breadth.
- 215 2. **MCI:** High value indicates the model tends to repeat answers across paraphrases, signaling
216 memorization.

217 3. **CRMA:**

$$Acc_{adj} = 0.85 \times (1 - 0.25 \times 0.70) = 0.85 \times 0.825 = 0.701$$

218 The adjusted accuracy falls to 70.1%, showing that contamination and memorization materi-
219 ally inflate reported performance.

220 **Example 2**

221 Let us consider evaluating a model on the DROP (Discrepancy in Reading Comprehension) bench-
222 mark, which is specifically designed to test reasoning over passages with arithmetic and logical
223 operations. Suppose the dataset has minimal overlap with pretraining corpora (low contamination),
224 and the model is observed to adapt flexibly across paraphrased question variants.

225 **Assumptions:**

- 226 • Raw accuracy (Acc): 72%
- 227 • Benchmark Ecology Score (BES): Medium Bias (divergence detected but with sufficient
228 topical coverage)
- 229 • Contamination probability ($\hat{C}$): 0.05 (low, as benchmark items are adversarially generated)
- 230 • Memorization Consistency Index (MCI): 0.25 (low, indicating reliance on reasoning rather
231 than rote recall)
- 232 • Function $f(MCI) = MCI$ (linear scaling)

233 **Inferences:**

- 234 1. **BES:** Medium Bias $\Rightarrow$ Benchmark is somewhat representative but not overly narrow.
- 235 2. **MCI:** Low value suggests model responses vary across paraphrases in meaningful ways,
236 pointing toward reasoning reliance rather than memorization.
- 237 3. **CRMA:**

$$Acc_{adj} = 0.72 \times (1 - 0.05 \times f(0.25)) = 0.72 \times (1 - 0.0125) = 0.72 \times 0.9875 = 0.711$$

238 Hence, the adjusted accuracy remains close to raw accuracy, reflecting that contamination is
239 not materially inflating the scores.

240 **Example 3**

241 Let us evaluate a model on a subset of MMLU (Massive Multitask Language Understanding), where
242 items range from definitional recall to light reasoning. The benchmark is reasonably aligned with
243 real-world knowledge queries (good ecology), but historical public availability of some items induces
244 moderate contamination. Model behavior suggests a mix of recall and reasoning.

245 **Assumptions:**

- 246 • Raw accuracy (Acc): 78%
- 247 • Benchmark Ecology Score (BES): Low Bias (close alignment to reference corpus)
- 248 • Calibrated contamination probability ($\hat{C}$): 0.15 (moderate)
- 249 • Memorization Consistency Index (MCI): 0.45 (medium; mixed signals)
- 250 • Scaling function: $f(MCI) = MCI$ (linear)

251 **Inferences:**

252 1. **BES:** Low Bias $\Rightarrow$ the benchmark is broadly representative; ecology alone would not explain
253 inflated scores.

254 2. **MCI:** Medium value indicates partial reliance on recall with some genuine generalization.

255 3. **CRMA:**

$$Acc_{\text{adj}} = 0.78 \times (1 - 0.15 \times f(0.45)) = 0.78 \times (1 - 0.0675) = 0.78 \times 0.9325 = 0.72735 \approx 72.7\%$$

256 The adjustment is noticeable (reflecting moderate contamination and mixed memorization)
257 but smaller than in high-contamination cases.

258 Therefore, the metrics of DC-Guard: BES contextualizes the benchmark scope, MCI highlights
259 possible memorization artifacts, and CRMA yields an adjusted, contamination-resilient performance
260 score that better reflects generalization.