

Static Benchmarks Are Not Enough: Why Frontier AI Evaluation Must Be Continuously Maintained

Version 1.0 · Published July 17, 2026 · Current · Evidence reviewed through July 17, 2026

This essay draws from Foundation 1, Dynamic Evaluation Protocols, within Standards Body's Foundations for Frontier AI research program. It does not replace Foundation 1 and does not constitute an adopted standard.

Approved Standards Body Public Research Publication. Publication status: Published. Canonical status: Noncanonical. Authority effect: None.

This publication is issued by Standards Body solely under its project-level authority over its own research and publications. Standards Body is an independent research and institutional-design project developing foundations for frontier AI evaluation, assurance, standards, and governance.

This publication is a public-facing derivative of Foundation 1, Dynamic Evaluation Protocols. It does not replace, amend, or supersede Foundation 1.

This publication is not: a technical standard, a draft standard, a certification criterion, an accreditation criterion, a regulatory requirement, governmental guidance, a legal conclusion, a safety determination, an evaluator approval, or a statement of industry-wide consensus.

Standards Body is not currently: a regulator, a governmental authority, a certification body, an accreditation body, a formally recognized standards-development organization, an evaluator approval body, or an operating assurance institution.

Opening

Static benchmarks have been essential to artificial intelligence research. They create stable reference points, support reproducible experiments, preserve historical baselines, enable comparison, and make technical progress easier to inspect. Those functions remain valuable.

The problem is not that static benchmarks are worthless. It is that they are often asked to support claims and decisions after their evidentiary conditions have changed.

Frontier AI systems develop quickly. Training and post-training methods change. Tools, retrieval, memory, scaffolding, and safeguards alter system performance. Public tasks circulate through papers, repositories, evaluation pipelines, training corpora, and developer workflows. Models may be optimized directly or indirectly against familiar measurements. New capabilities, deployment settings, and failure modes appear that older task sets were not designed to represent.

A benchmark can keep producing a precise score while the inference behind it becomes narrower, less current, or less decision-relevant.

A measure may remain repeatable while losing validity for the interpretation or decision it is being used to support.

The central claim of this essay is bounded: static benchmarks remain useful, but they are insufficient as the sole basis for consequential frontier AI evaluation. Such evaluation should be organized through versioned, governed, and continuously maintained protocols.

The concept is **dynamic evaluation**: continuously maintained, evidence-responsive evaluation governed through controlled change. It does not mean changing tests constantly or moving thresholds when results are inconvenient. It means maintaining evaluation infrastructure so that the evidence remains worthy of reliance.

A Benchmark Is Not a Complete Evaluation

A benchmark is a standardized comparison instrument. An evaluation protocol is the complete specification governing an evaluation.

A benchmark score does not, by itself, identify everything needed for a consequential interpretation. It may not establish:

- The exact model or system evaluated
- The prompts, tools, retrieval, memory, safeguards, or human assistance available
- The administration and elicitation conditions
- Whether tasks were exposed, contaminated, or compromised
- Whether performance was optimized against the measure
- How uncertainty was estimated
- Whether the result generalizes beyond the tested conditions
- Which decision the result is intended to inform
- How long the result should remain current

For frontier AI, a model name is often an inadequate description of the evaluated object. A model and a deployed system are not interchangeable. The same model can behave differently with different tools, prompts, access controls, monitoring, or human oversight.

Evaluation conditions are therefore part of the result, not incidental metadata.

A credible protocol defines the question, evaluated object, construct, tasks, administration, elicitation, scoring, uncertainty, integrity controls, reporting, version, and intended use. A benchmark may be one component. It should not be mistaken for the whole.

This distinction also limits what evaluation can establish. Capability and risk are not interchangeable. Capability evidence may inform risk analysis, but risk also depends on actors, access, exposure, safeguards, context, likelihood, consequence, and uncertainty. Passing an evaluation does not prove that a system is safe. Failure to demonstrate a capability under assessed conditions does not prove that the capability is absent.

Why Fixed Benchmarks Lose Decision Value

Fixed benchmarks can lose decision value through several mechanisms.

Saturation occurs when top systems cluster near a benchmark’s ceiling or when score differences become too small to distinguish meaningful capability differences. Small differences may then reflect sampling noise, prompting choices, or scoring details more than meaningful differences in the capability of interest. Making a test harder does not automatically restore validity. Difficulty is not the same as representativeness or decision relevance.

Contamination and exposure weaken the inference that performance reflects general capability rather than familiarity, memorization, or targeted preparation. Exposure does not make every result useless, but it changes what the result can reasonably support.

Optimization against the measure creates a related problem. Once a benchmark becomes influential, developers have incentives to improve performance on it. Some improvement reflects genuine capability. Some may reflect benchmark-specific tuning or narrower adaptation. The score alone may not reveal the difference.

A recent coding evaluation makes this concrete. LiveCodeBench compared HumanEval+ with the LCB-Easy code-generation subset for problems released from September 2023 to May 2024. The paper reported a correlation of 0.72 and marked model-level differences. DS-INS-1.3B, for example, received Pass@1 scores of 60 on HumanEval+ and 26 on LCB-Easy. The authors described the broader pattern as indicating potential overfitting on HumanEval+, especially among several fine-tuned open-access models. This does not make HumanEval+ invalid or useless, and the paper does not establish contamination as the cause of every difference. It shows that strong performance on a familiar benchmark may transfer imperfectly to newer tasks and may support a narrower inference than continued use of the familiar score suggests.

Changes in system boundaries further reduce comparability. Tools, retrieval, scaffolds, safeguards, and user interaction may materially alter behavior. Comparisons become unreliable or invalid when materially different configurations are treated as equivalent.

Construct drift occurs when the question that matters changes while the benchmark remains fixed. A test built around short, self-contained tasks may provide weak evidence about long-horizon agents, tool-using systems, recovery from failure, strategic behavior, or deployment under uncertainty. The benchmark may still measure its original tasks accurately while becoming less relevant to the decision now under consideration.

These mechanisms may leave the benchmark’s familiar name, interface, and numerical output intact. The deeper risk is expired or narrowed evidence presented in a current-looking form.

What Dynamic Evaluation Actually Means

Dynamic evaluation treats evaluation as a maintained evidence system, not a finished dataset.

A dynamic protocol should identify which elements must remain stable and which may change in response to evidence. Stable elements may include the underlying construct, intended decision, core definitions, minimum system-identity requirements, and governance principles. Dynamic elements may include task samples, adversarial methods, administration environments, scoring details, elicitation budgets, tools, security controls, monitoring triggers, and expiration rules.

Changes should follow defined triggers, such as task compromise, saturation, material system changes, improved elicitation, new deployment contexts, incidents, emerging failure modes, or evidence that the instrument no longer represents the outcome of interest.

LiveBench illustrates how a benchmark can be actively maintained rather than published once and left unchanged. Its authors designed it around frequently refreshed questions and the addition of new and harder tasks as systems improve. Its implementation also supports release-specific evaluation, allowing results to be associated with dated benchmark releases. The repository documents a dated release for which not all questions were simultaneously public on Hugging Face. LiveBench remains a maintained benchmark or evaluation instrument, not a complete evaluation protocol. Its design nevertheless demonstrates why content renewal should be paired with explicit release management. Refresh without versioning can make historical comparison less interpretable. Maintenance can limit contamination risk and preserve discrimination, but it does not guarantee validity or permanent freedom from contamination.

Dynamic evaluation should also be proportionate. A public research comparison and a high-consequence deployment decision do not require identical governance. As error consequences and public reliance increase, so should demands for complementary evidence, independent challenge, justified protection, and documented maintenance.

What Credible Maintenance Requires

Credible maintenance rests on four connected ideas.

1. Protocol and Evaluated-System Identity

Every consequential result should identify the protocol version, evaluated system, material configuration, date, conditions, evaluator, and relevant tool or access settings. Results should not transfer automatically across versions, configurations, or deployments.

2. Governed Change and Version Control

Protocol changes should be documented and reviewable. Maintainers should state what changed, why, which evidence triggered revision, who authorized it, what validation occurred, and whether earlier results remain current.

Revision should preserve history rather than overwrite it. Release identifiers, change logs, archived task definitions where safe, and visible status labels make it possible to distinguish correction, refresh, supersession, and retirement.

3. Comparability, Uncertainty, and Expiration

A new protocol release does not automatically produce results directly comparable with earlier releases. Maintainers may use anchor tasks, overlapping samples, reference systems, calibration, or bridge studies to assess continuity.

Where comparability can be defended, the basis should be explained. Where it cannot, discontinuity should be explicit. Sometimes the responsible conclusion is that a new baseline is required.

Results should also be time-bounded. A consequential report should identify uncertainty, a valid-through date or review condition, and reevaluation triggers. One configuration's result should not become a permanent property of a model family.

4. Independent Review, Correction, and Retirement

Dynamic evaluation creates discretion, and discretion requires challenge. Independent reviewers should be able to examine the protocol's purpose, methods, evidence, limitations, conflicts, and change history, with the depth of review proportionate to the decision.

Tasks may sometimes need protection to preserve integrity or prevent harmful disclosure. That can reduce public reproducibility. Maintained evaluation also raises infrastructure costs and can make comparisons across releases more difficult, increase fragmentation among incompatible systems, expand maintainer discretion, and create risks of capture or selective revision.

Operational feasibility also depends on secure task custody, reproducible execution environments, qualified evaluator access, release engineering, and funding sufficient to conduct refresh and bridge studies without lowering review quality.

These tradeoffs are real. They strengthen the case for release identifiers, transparent governance, documented change procedures, independent review, and explicit discontinuity. They do not support leaving consequential evaluation static by default.

Correction and retirement are part of evaluation quality. Tasks can be compromised, scoring can fail, and constructs can change. A protocol should be able to narrow claims, correct results, retire components, or withdraw interpretive authority without erasing history.

Governance Is Part of Measurement

Dynamic evaluation is both technical and institutional.

Someone must decide what remains stable, which evidence triggers revision, how uncertainty is reported, when comparison has failed, when a result expires, and when a protocol should be retired. These choices can influence deployment, procurement, research funding, insurance, standards, or regulation. They should not remain hidden inside an evaluation pipeline.

A developer-issued frontier framework already reflects this iterative logic. Google DeepMind's Frontier Safety Framework 2.0 sets out regular evaluation of its most powerful frontier models, additional evaluation when a model may represent an exceptional capability increase, and more frequent early-warning evaluation or adjusted alert thresholds when the safety buffer may be insufficient. It also provides for periodic review of the appropriateness and efficacy of applied mitigations, continued post-mitigation testing, updated threat modeling and risk-landscape analysis, and post-deployment safety-case revision through red-teaming and threat-model revisions. The framework states that safeguards may also be updated to ensure continued adequacy. This is an official developer-issued framework, not a governmental requirement, public standard, regulatory obligation, industry consensus, universally adopted framework, or independently validated framework. It nevertheless shows how frontier evaluation can be treated as an iterative control process rather than a one-time release gate.

Credible governance requires clear roles, relevant competence, conflict management, documentation, and paths for independent challenge and correction. A technically sophisticated protocol may still lose credibility if revision authority is opaque, financially dependent, insulated from criticism, or controlled by actors with a direct interest in the outcome.

This essay advances five bounded recommendations for consequential frontier AI evaluation:

- 1 Treat the protocol, not merely the benchmark, as the governed unit.
- 2 Use controlled versioning and evidence-triggered maintenance.
- 3 Identify the evaluated system, conditions, date, uncertainty, and expiration.
- 4 Preserve comparability where defensible and disclose discontinuity where it is not.
- 5 Apply stronger governance and complementary evidence as the consequences of error increase.

These research-based recommendations are not an adopted standard, requirement, certification criterion, accreditation criterion, regulatory rule, or legal conclusion.

Conclusion

Static benchmarks remain useful as stable reference points, reproducible research instruments, historical baselines, comparison tools, and components within broader evidence portfolios.

They are not enough on their own.

A fixed task set cannot ensure that evidence remains current as models, systems, environments, threats, and decision needs change. Consequential frontier AI evaluation should therefore move from benchmark-centered thinking toward maintained protocol governance.

That means controlled change rather than arbitrary change. It means versioning rather than silent replacement. It means preserving comparison where evidence supports it and disclosing discontinuity where it does not. It means identifying the evaluated system and conditions, attaching uncertainty and time to results, and permitting independent challenge, correction, and retirement.

It also means preserving the limits of evaluation. A well-governed dynamic protocol can support stronger bounded claims. It cannot prove that an AI system is universally safe, harmless, compliant, or suitable for every deployment.

Three questions remain especially important:

- 1 How should benchmark decay and compromise be detected?
- 2 How should comparability be preserved, or honestly discontinued, when protocols change?
- 3 Who should govern protocol revision, challenge, expiration, and retirement?

Those questions require further empirical and institutional work. Frontier AI evaluation should not depend solely on measurements that remain fixed while the systems and decisions around them change.

It should be continuously maintained because the value of evaluation lies not in producing a permanent score, but in producing evidence that remains worthy of reliance.

Source Notes

- 1 [Naman Jain et al., “LiveCodeBench: Holistic and Contamination Free Evaluation of Large Language Models for Code,”](#) ICLR 2025, Section 5.2 and Figure 4. The paper reports a 0.72 correlation between HumanEval+ and the LCB-Easy code-generation scenario for the September 2023 to May 2024 window. It identifies DS-INS-1.3B as an example with Pass@1 scores of 60 and 26 and interprets the broader pattern as potential overfitting on HumanEval+, particularly among several fine-tuned open-access models.
- 1 [Colin White et al., “LiveBench: A Challenging, Contamination-Limited LLM Benchmark,”](#) ICLR 2025 Spotlight, together with the [official LiveBench GitHub repository](#), including its README and changelog, accessed July 17, 2026. The paper describes frequently updated questions, recurring additions, and new or harder tasks over time. The repository implements dated release options and has documented that not all questions for one dated release were simultaneously public on Hugging Face.
- 1 [Google DeepMind, “Frontier Safety Framework,” Version 2.0,](#) February 4, 2025, Sections 2 and 3. The framework describes regular and triggered frontier-model evaluation, early-warning evaluation adjustments, periodic mitigation review, continued post-mitigation testing, updated threat modeling, post-deployment safety-case revision through red-teaming, and possible safeguard updates.

Citation: Standards Body. "Static Benchmarks Are Not Enough: Why Frontier AI Evaluation Must Be Continuously Maintained." Public Foundational Essay SB-PUB-2026-0001, Version 1.0, July 17, 2026.

<https://standardsbody.ai/library/publications/static-benchmarks-are-not-enough/>

Immutable Version 1.0: <https://standardsbody.ai/library/publications/static-benchmarks-are-not-enough/versions/1.0/>

Corrections: contact@standardsbody.ai. Corrections, supersession, and withdrawal are recorded; Version 1.0 is preserved without silent edits.